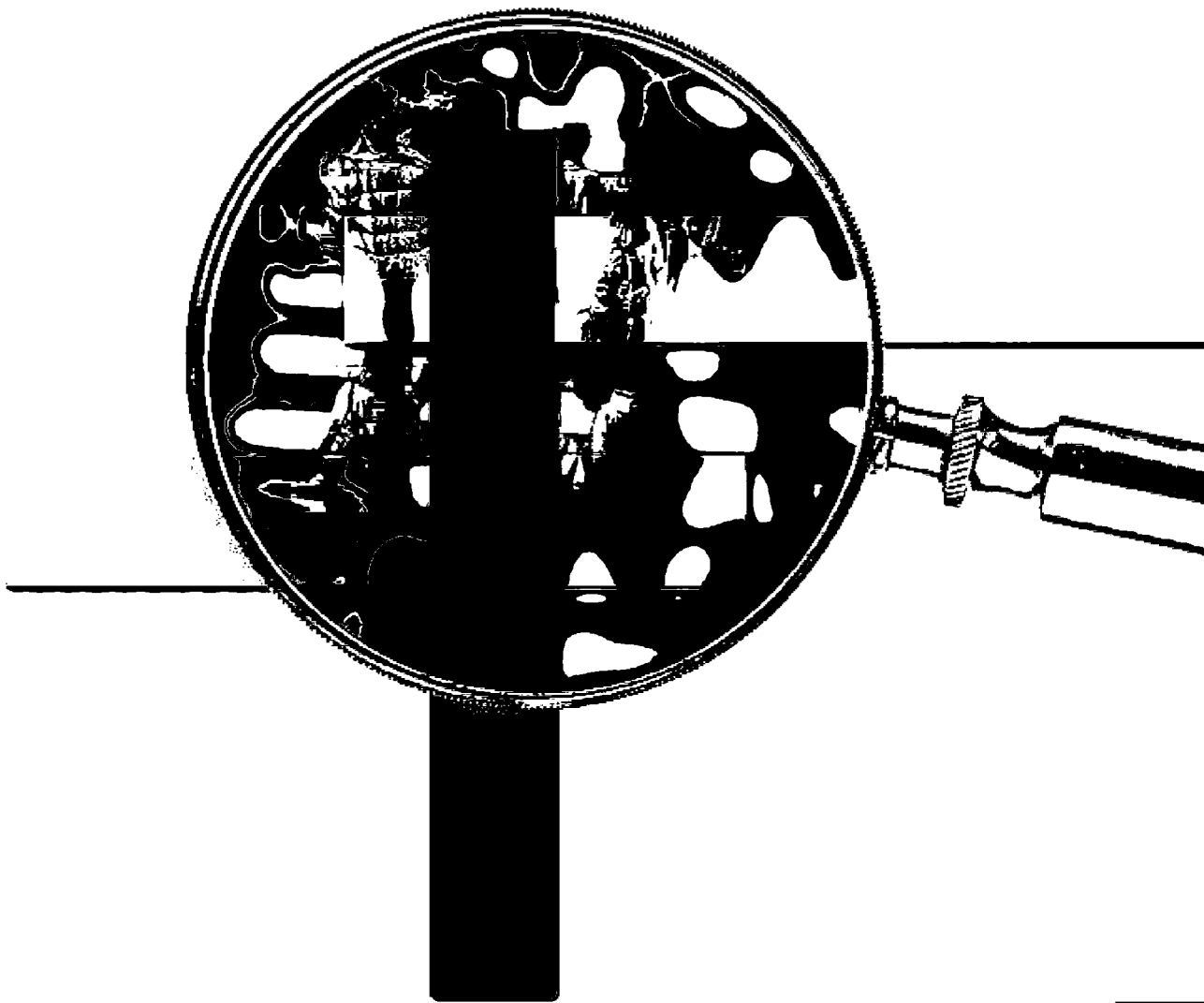


ganz1912

Técnicas de investigación aplicadas a las ciencias sociales

JORGE PADUA
(coordinador)



SOCIOLOGÍA

Jorge Padua nació en Argentina en 1938; estudió psicología en la Universidad Nacional de Tucumán (Argentina); obtuvo el grado de maestro en la Escuela Latinoamericana de Sociología (ELAS), en Chile, y el de doctor en la Universidad de Alberta, en Canadá. Desde 1973 fue profesor e investigador en El Colegio de México y luego de su jubilación fue recontratado por la misma institución para el Programa de Estudios Interdisciplinarios del Centro de Estudios Sociológicos.

Ingvar Ahman fungió como experto de la UNESCO en la Escuela Latinoamericana de Sociología (ELAS), de la Facultad Latinoamericana de Ciencias Sociales (FLACSO), en Chile. Jubilado de la Organización Mundial de la Salud, reside actualmente en Suiza.

Héctor Apezechea fue un sociólogo uruguayo egresado de la Facultad Latinoamericana de Ciencias Sociales (FLACSO), en Chile. Se desempeñó como consultor de la Organización Panamericana de la Salud y como investigador del Centro de Informaciones y Estudios del Uruguay (CIESU).

Carlos Borsotti es maestro en sociología por la Escuela Latinoamericana de Sociología (ELAS) y doctor en ciencias jurídicas y sociales por la Universidad Nacional del Litoral, de la que también es profesor de posgrado en ciencias sociales. Asimismo, es profesor titular ordinario en la Universidad Nacional de Luján.

SECCIÓN DE OBRAS DE SOCIOLOGÍA

TÉCNICAS DE INVESTIGACIÓN APLICADAS
A LAS CIENCIAS SOCIALES

TÉCNICAS DE INVESTIGACIÓN APLICADAS A LAS CIENCIAS SOCIALES

JORGE PADUA

INGVAR AHMAN

HÍCTOR APEZECHEA

CARLOS BORSOTTI



FONDO DE CULTURA ECONÓMICA
EL COLEGIO DE MÉXICO

Primera edición, 1979

Primera edición electrónica, 2018

D. R. © 1979, Fondo de Cultura Económica

Carretera Picacho-Ajusco, 227; 14738 Ciudad de México



www.fondodeculturaeconomica.com

Comentarios:

editorial@fondodeculturaeconomica.com

Tel. (55) 5227-4672

Diseño de portada: Laura Esponda Aguilar

Se prohíbe la reproducción total o parcial de esta obra, sea cual fuere el medio. Todos los contenidos que se incluyen tales como características tipográficas y de diagramación, textos, gráficos, logotipos, iconos, imágenes, etc., son propiedad exclusiva del Fondo de Cultura Económica y están protegidos por las leyes mexicanas e internacionales del copyright o derecho de autor.

ISBN 978-607-16-5025-2 (mobi)

Hecho en México - *Made in Mexico*

ganz1912

PRÓLOGO

El texto que presentamos al lector no es un manual de metodología, sino de técnicas de investigación y, dentro de esta área, centra el interés en los métodos de *survey* o investigaciones por cuestionarios (preferimos la expresión *survey* en la medida en que la palabra evoca el proceso de investigación antes que el instrumento). El manual está escrito con el supuesto general de que las ciencias sociales son ciencias empíricas y nomotéticas, cuyo objetivo general es describir, predecir y explicar.

En el texto no nos interesa tanto la investigación normativa como la investigación descriptiva: exponemos al lector una serie de técnicas e instrumentos, señalando su modo de construcción, sin profundizar en la racionalidad de los mismos. En la medida en que el texto es de carácter introductorio, constantemente invitamos al lector a profundizar en los temas, vía literatura especializada en el área y para la cual se proporcionan algunas referencias bibliográficas.

La peculiaridad o característica esencial del tipo de investigación que trata el manual está definida por el hecho de que la recolección sistemática de datos se establece a través de contactos directos o indirectos con individuos o grupos de individuos, es decir, con algún tipo de informante. La encuesta y las pautas de investigación sobre las que trata el manual son procedimientos comúnmente utilizados en la investigación social. Existen por supuesto numerosas y distintas técnicas y procedimientos, ni menos ni más relevantes a la investigación. Los métodos y las técnicas no se definen como apropiadas o no apropiadas *a priori*, sino que dependen del problema que la investigación pretende resolver, del estado de avance de la teoría sustantiva y del tipo de pregunta que el investigador tratará de responder.

La problemática teórica define tanto el objeto como los métodos con que se apropia el objeto. Las técnicas de recolección de datos, así como la estadística, son instrumentos de los cuales se puede servir el investigador, dados sus propósitos teóricos. De ahí que deba tomar frente a ellos una actitud a la vez flexible y vigilante. Flexible en el sentido de no juzgarlos *a priori*; crítica y vigilante en el sentido de no utilizarlos automáticamente.

Pero el manual que sometemos al lector, repetimos, no es un manual de metodología y menos de epistemología, sino un manual técnico, si se quiere una especie de “recetario” en el que se tratan de señalar y especificar en su mayor amplitud los aspectos relativos a la parte operacional de la investigación por encuestas. Partimos de la base de que el investigador ha resuelto parcialmente sus problemas teóricos y tiene clara conciencia de que, por ejemplo, no tiene sentido aplicar un muestreo aleatorio simple en investigaciones que tienen como objeto determinar relaciones estructurales o procesos de interacción, ya que por las características mismas del muestreo aleatorio simple se destruyen las redes de relaciones que se pretende estudiar.

El texto comenzó a escribirse en Santiago de Chile, en la Escuela Latinoamericana de Sociología, en los años 1966 y 1967, con Ingvar Ahman, entonces experto de la UNESCO en la ELAS. Posteriormente Ahman viajó a Suecia y nosotros a Canadá. En nuestro reencuentro en México, nos encargamos de replantear la estrategia del manual, quedando Padua como principal responsable. Algunos de los capítulos fueron totalmente reconstruidos y otros parcialmente modificados. Decidimos conservar de la versión original íntegramente los capítulos de Héctor Apezechea sobre “Codificación” y de Carlos Borsotti sobre “Análisis de datos: el concepto de propiedad-espacio y la utilización de razones, tasas, proporciones y porcentajes”.

El capítulo I, escrito por Ingvar Ahman y Jorge Padua, es de carácter general, y sintetiza el manual a través de un diagrama sobre la organización de un *survey* que contiene seis pasos: 1) Orientación en el campo de investigación; 2) Construcción, evaluación y manejo del instrumento de recolección de datos; 3) La recolección de datos; 4) El procesamiento; 5) El análisis y la interpretación, y 6) La presentación.

Los capítulos subsiguientes toman algunos de los casilleros que aparecen en el diagrama de la página 10 y los desarrollan con detalle.

El capítulo II trata del proceso de investigación, concentrando mayormente su interés en la *operacionalización* de variables y la construcción de índices simples. Fue escrito por Jorge Padua.

El capítulo III trata del muestreo, delineándose algunas indicaciones sobre la construcción de muestras probabilísticas y muestras no probabilísticas. Se contempla además la idea de Galtung, en su libro *Teoría y métodos de la investigación social* (Eudeba, Buenos Aires, 1966), de la construcción de muestras para la prueba de hipótesis sustantivas. El capítulo fue escrito por Jorge Padua.

El capítulo IV, de Jorge Padua e Ingvar Ahman, trata sobre el cuestionario, especialmente en lo referente a los aspectos técnicos en la construcción de los mismos, ejemplificando alternativas para la formulación de preguntas y en su ordenamiento.

El capítulo V, escrito por Héctor Apezechea, está referido a los procedimientos de codificación, confección de código y procedimientos de revisión y control.

El capítulo VI, escrito por Jorge Padua e Ingvar Ahman, presenta una serie de escalas para la medición de actitudes, señalando para la mayoría de ellas las técnicas para su construcción y comparándolas en términos de su eficiencia como instrumentos de medición.

El capítulo VII, escrito por Ingvar Ahman sobre trabajo de campo en investigaciones por cuestionarios, particularmente cuando éstas son realizadas a gran escala, propone criterios administrativos importantes para la organización y el buen éxito de la recolección del material.

El capítulo VIII, escrito por Carlos Borsotti, cubre la primera parte del análisis de datos y abarca dos áreas importantes: el concepto de propiedad-espacio y algunas estadísticas de nivel nominal como razones, proporciones, tasas y porcentajes. El capítulo abunda en ejemplos y es una contribución importante a la idea original de A. Barton (“The Concept of Property Space in Social Research”, en Lazarsfeld, P., y

Rosenberg, M.: *The Language of Social Research*; The Free Press of Glencoe, Illinois, 1955).

El capítulo IX, escrito por Jorge Padua, cubre otra parte de la sección dedicada al análisis estadístico de los datos. Dada la abundancia en el mercado de textos de estadística, hemos preferido integrar el uso de computadoras y la inclusión de paquetes estadísticos en ellas, para señalar la oferta y las condiciones para el uso e interpretación de un programa: el SPSS (paquete estadístico para las ciencias sociales). Desarrollamos con mayor detalle en este capítulo la parte correspondiente a los métodos más refinados de análisis (análisis de la varianza, análisis factorial, análisis discriminante, análisis del escalograma Guttman), dedicando escasa atención a la estadística descriptiva e inferencial, material y temas cuyo desarrollo es relativamente más fácil de encontrar en el medio. El capítulo está basado principalmente en el texto de Nie, N., Hull, C. H., y Jenkins, J. *Statistical Package for the Social Sciences*; McGraw-Hill, Nueva York, 1975 (2ª ed.).

El último capítulo, escrito por Ingvar Ahman, se refiere a algunas recomendaciones generales sobre la presentación del informe de investigación, señalando algunos criterios acerca de lo que debe ser incluido en el mismo, con el fin de facilitar la comunicación con la comunidad académica y de investigadores.

Todos los capítulos han sido escritos de manera tal que puedan ser leídos en forma independiente, de ahí que algunos temas aparezcan repetidos a lo largo del libro. Sin embargo, su ordenamiento en el texto sigue la lógica del proceso de investigación.

Las enseñanzas de Ingvar Ahman en la Escuela Latinoamericana de Sociología (ELAS-FLACSO), así como su entusiasmo en la idea de publicar un texto accesible en castellano, han sido los principales “motivos” para publicar el manual. Desafortunadamente responsabilidades con UNDP le impidieron concentrar sus esfuerzos para que el manual resultara una responsabilidad editorial compartida. De todas formas, es a Ahman a quien expresamos nuestro mayor agradecimiento por su contribución teórica, su participación en varios capítulos y su insistencia en buscar un nivel que resultara lo más accesible al lector no familiarizado con las matemáticas o con la “jerga” técnica.

Agradecemos a Claudio Stern por sus útiles comentarios y observaciones críticas al borrador de este manual. Nos hemos beneficiado asimismo de los trabajos y observaciones críticas de las promociones V y VI de la Escuela Latinoamericana de Sociología, principalmente en la primera versión del manual.

El Centro de Estudios Sociológicos de El Colegio de México y su director, Rodolfo Stavenhagen, nos brindaron todo el aliento para emprender la tarea y costearon parte del tiempo, permitiendo además que distrajáramos parcialmente nuestra atención de la tarea docente y de investigación.

Guadalupe Luna transcribió el texto a máquina y tuvo que sufrir las atrocidades de mi caligrafía y ortografía.

Todas esas personas e instituciones han desempeñado un papel importante en la gestación y carácter del manual y a ellos toda mi gratitud. Sin embargo, la propiedad de los errores sigue siendo del dominio exclusivo del autor responsable, confiando en que

ellos sean tomados con benevolencia.

Esperamos que el manual tenga utilidad en el terreno para el cual está destinado: los cursos introductorios de técnica de la investigación y los investigadores que necesitan a menudo consultar sobre aspectos técnicos de la investigación, asuntos que desafortunadamente no abundan en la bibliografía especializada.

JORGE PADUA

I. LA ORGANIZACIÓN DE UN *SURVEY*

JORGE PADUA
INGVAR AHMAN

EL OBJETIVO de este capítulo es proporcionar una breve y esquemática introducción acerca de la idea de cómo llevar a cabo una investigación de tipo *survey*. Con el fin de clarificar el proceso, usaremos un diagrama en el que señalamos las distintas pautas que componen el proceso de investigación. El diagrama no pretende incluir, naturalmente, todas las variaciones y diferentes pasos que debe seguir una investigación, sino más bien nos servirá de pauta. El campo de las investigaciones tipo *survey* incluye varias técnicas de recolección de datos, pero aquí sólo trataremos en detalle la más utilizada: el cuestionario. Otras técnicas, como por ejemplo las observaciones participantes, las entrevistas clínicas (en lo que se refiere a datos primarios) y el análisis de contenido y otras técnicas (en lo que se refiere a datos secundarios), tienen importancia, pero serán tratadas secundariamente en otras secciones de este libro. Los pasos en el diagrama son:

- Paso I: Orientación en el campo de investigación y formulación de un sistema de hipótesis.
- Paso II: La construcción, evaluación y manejo del instrumento de recolección de datos (cuestionario) y muestreo.
- Paso III: Recolección de datos.
- Paso IV: El procesamiento de los datos.
- Paso V: El análisis.
- Paso VI: La presentación.

Trataremos cada uno de ellos por separado.

PASO I: ORIENTACIÓN EN EL CAMPO DE INVESTIGACIÓN Y FORMULACIÓN DE UN SISTEMA DE HIPÓTESIS

El primer paso que debe dar el investigador —no importa por quién esté patrocinada la investigación o qué motivos lo hayan impulsado a efectuar el estudio— es tener una sólida orientación en el campo que va a investigar. Esta orientación se refiere a las elaboraciones abstractas de la teoría, a los resultados de investigación y a las particulares circunstancias concretas que constituyen el objeto o situación a investigar.

Documentación descriptiva:

La literatura actual y los documentos históricos de información pueden dar luz a los problemas que investigará, sobre todo en relación a los aspectos y peculiaridades concretas. Una revisión de informaciones de prensa, radio y televisión puede resultar también de mucha utilidad, sobre todo cuando se trata de hacer un análisis de contenido en investigaciones donde el objetivo es la medición de actitudes u opiniones, por ejemplo.

La consulta de archivos públicos y de documentos oficiales es también útil. Si existe el interés o la necesidad de extraer de ellos algunos datos, el trabajo es más arduo, pues depende en buena medida de la organización interna del archivo.¹

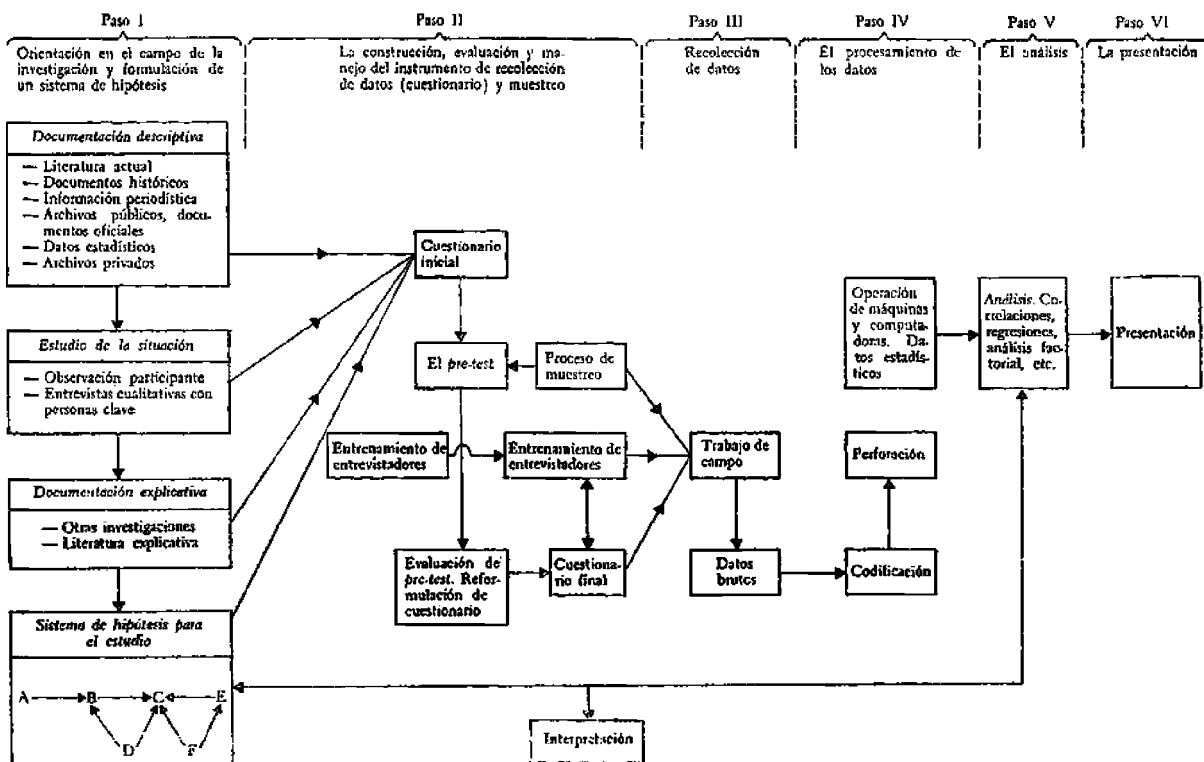
El estudio puede también recurrir a las informaciones disponibles en archivos privados, en los cuales hasta los memoranda y las notas escritas a mano pueden ser usados. Debe tenerse especial cuidado con el material estadístico disponible, el cual debe ser estudiado detenidamente, ya que la utilidad de este material depende de la manera en que fue obtenido y calculado.²

Estudio de la situación

En muchos casos, y si el área de investigación es totalmente desconocida para el investigador, es recomendable un primer contacto como observador participante. El investigador debe, por ejemplo, vivir con la gente que quiere estudiar; tomar un trabajo en una fábrica del mismo tipo que la que va a estudiar, o actuar como espectador en un determinado ambiente.³

Como complemento a los contactos como observador participante, es de gran importancia llevar a cabo algunas entrevistas no estructuradas, con algunos sujetos llamados “personas clave”. En estas entrevistas hay que buscar que el sujeto entrevistado esté en su ambiente natural, de tal modo que pueda contar sus deseos y opiniones con referencia a los problemas en los cuales el investigador está interesado, actuando de la manera más natural posible. De este modo el resultado de la entrevista es bastante más confiable. La persona “clave” puede ser alguien que ocupe una situación destacada dentro de una industria, el presidente de un sindicato, algunos periodistas especializados en la materia, un político que ha usado el problema en su campaña o labor, etc. La idea es utilizar las entrevistas no estructuradas y hacerse un cuadro amplio de la situación y de los problemas involucrados, a través de lo que diferentes personas piensan o actúan con respecto a estos problemas y de las actitudes que podrían tomar.

Diagrama del proceso de investigación para un diseño tipo survey, con cuestionario como instrumento para la recolección de datos



Estas entrevistas deben ser realizadas por el mismo investigador. El número de entrevistas varía de acuerdo con el tamaño e importancia del *survey* y del campo que el investigador estudiará. Habrá casos en que será suficiente con cinco entrevistados, y otros en que se necesitarán 30 o más. Es recomendable que, aunque la entrevista se haya descrito como no estandarizada, el entrevistador agregue algunos puntos; es decir, que fije algunas áreas o preguntas que desee cubrir y entonces espere la reacción del entrevistado. Si el sujeto abre nuevos campos de interés, el entrevistador debe seguir estos campos, desarrollándolos.

Con estas entrevistas, un aspecto *cualitativo* entra en la investigación. Queremos acentuar la importancia de esta manera de trabajar, pues luego de la exploración es posible comenzar a formalizar la hipótesis o incluso cambiar todo el carácter de la investigación. El proceso que sigue todo el primer paso de orientación de una materia y la formalización de una hipótesis se concentra en la recolección de datos a través de un cuestionario y el análisis estadístico que da a la investigación su carácter principal de ser *cuantitativa*.

Documentación explicativa

Aquí se trata de analizar qué han pensado y expresado sobre el problema que nos interesa otros autores, cómo han afrontado y formulado el problema, cómo lo han resuelto, a qué conclusiones han llegado, cómo han definido sus conceptos, cómo han determinado sus observaciones, etcétera.

Las entrevistas cualitativas son el primer intento para integrar la documentación

descriptiva y la conceptualización más o menos generalizada a la situación concreta de la investigación particular. A este nivel el investigador comienza a definir sus preguntas más específicamente, así como la formulación de sus primeras hipótesis.

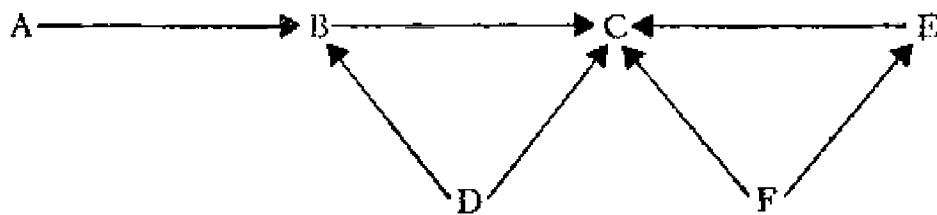
Ahora bien, el investigador debe llevar el sistema al papel. Como ayuda puede consultar los resultados de otras investigaciones efectuadas en la misma área. Habrá en ellas seguramente algo escrito sobre las variables utilizadas y cómo fueron construidas e interconectadas y a qué conclusiones llegaron. Es muy importante consultar la literatura explicativa o teórica. A través de ella el investigador puede insertar su estudio particular en un marco de referencia teórica más general.

En algunos casos puede suceder que el investigador prefiera dar a su investigación la forma de una réplica a otro estudio ya efectuado o aplicar una parte de un sistema teórico y tratar de establecer su aplicabilidad a otra área de la realidad. Este modo de utilizar ideas no tiene nada que ver con un “plagio”: por el contrario, es la manera en que la ciencia opera con mayor frecuencia con objeto de avanzar y ampliar el campo de la teoría. Por este motivo, en el último capítulo de este libro hemos hecho un llamado a los investigadores, pidiéndoles que presenten sus instrumentos (cuestionarios) y hallazgos en una forma sistematizada, con objeto de que otros investigadores puedan seguir ensayando los resultados una y otra vez.

Resumen

Como un resumen del primer paso de la manera de efectuar una investigación *survey*, recomendamos lo siguiente:

- 1) Documentación descriptiva: literatura actual; documentos históricos; revisión de informaciones de prensa, radio y T. V.; archivos públicos, documentos oficiales, datos estadísticos; archivos privados, documentos privados, datos estadísticos.
- 2) Estudio de la situación: observación participante; entrevistas cualitativas con personas clave.
- 3) Documentación explicativa: otras investigaciones; literatura explicativa.
- 4) Sistema de hipótesis para el estudio. Ejemplo:



A y F son variables; —————> indica una relación y su dirección

PASO II: LA CONSTRUCCIÓN, EVALUACIÓN Y MANEJO DEL INSTRUMENTO PARA LA RECOLECCIÓN DE DATOS (CUESTIONARIO) Y MUESTREO

En este paso de la investigación se trata de construir un instrumento que sirva para medir los conceptos que hemos seleccionado. Los métodos de recolección más utilizados en el tipo de investigación que estamos desarrollando son la observación, la entrevista y el cuestionario. Luego de un breve resumen descriptivo sobre cada uno de estos métodos, concentraremos nuestra atención en el cuestionario como instrumento de recolección de datos.

Observaciones, entrevistas y cuestionarios

La *observación* —como método de recolección de datos— se aplica preferentemente

en aquellas situaciones en las que se trata de detectar aspectos conductuales, como ocurre en situaciones externas y observables. Los cuestionarios y entrevistas se dirigen por lo general hacia la obtención de datos no observables directamente, datos que se basan por lo general en declaraciones verbales de los sujetos.

La *observación participante*, utilizada con gran eficacia por la antropología social, es especialmente indicada para propósitos exploratorios, y como señalábamos en los párrafos anteriores, forma parte del proceso de familiarización del investigador en el estudio de la situación. Aquí el análisis de los datos es simultáneo a la recolección de los mismos. El investigador tiene que determinar qué es lo que debe observar y cómo va a registrar esas observaciones. Lo que va a observar depende de la teoría en particular (implícita o explícitamente formulada). El registro de las observaciones debe hacerse tratando de minimizar el error en el registro, al mismo tiempo que evitar distorsionar la situación de observación. Por ejemplo, para evitar al máximo los errores de registro, productos de distorsiones en la memoria, puede ser conveniente tomar notas, utilizar grabadoras, filmaciones, etc. Al mismo tiempo, la utilización de estos mecanismos puede llevar a una perturbación en la situación, de manera que se pierda la espontaneidad de la misma. Hay que buscar —de acuerdo con cada situación— las soluciones que al incrementar la pureza del registro reduzcan la distorsión de la situación. Esto es posible mediante el uso de recursos mnemónicos o técnicas similares. Lo que interesa destacar aquí es que el investigador debe planear su estrategia anticipadamente, así como establecer listas y registros de observación de manera que la observación sea selectiva, concentrándose ésta en los detalles sustantivos.

La *observación sistemática* es bastante más estructurada que la observación participante; se aplica en situaciones de diagnóstico y clasificación en base a taxonomías o tipologías ya establecidas, de manera que las categorías de observación ya están codificadas, implicando la observación sistemática una tarea de registro bastante menos flexible que la de la situación en observaciones participantes. La observación en situaciones experimentales es estructurada ya a nivel de estandarización, con el fin de elevar el grado de comparabilidad de situaciones similares.

En el contexto de las investigaciones tipo *survey*, la *entrevista* es una técnica de recolección de datos que implica una pauta de interacción verbal, inmediata y personal, entre un entrevistador y un respondente. Las pautas de interacción entre entrevistador (*E*) y respondente (*R*) incluyen factores más complejos que el simple intercambio de estímulos y respuestas verbales. De esta manera la dinámica de la situación pasa de tratar la entrevista como un simple proceso mecánico de recolección de datos a una teoría psicológica de la situación de la entrevista, en la cual —como afirma Hyman—⁴ se postula la necesidad de tratar las respuestas individuales como síntomas, más que como realidades o hechos.

Dependiendo del tipo de investigación, las entrevistas se clasifican en estandarizadas, semiestandarizadas y no estandarizadas.

Las *entrevistas no estandarizadas* se utilizan en etapas exploratorias de la investigación, ya sea para detectar las dimensiones más relevantes, para determinar las

peculiaridades de una situación específica o para generar hipótesis iniciales. El rasgo esencial de este tipo de entrevistas es la flexibilidad en la relación entrevistador-respondente, lo que permite un margen tanto en la reformulación de preguntas como en la profundización en algunos temas y, por lo general, la ruptura en cualquier orden en cuanto a la secuencia en que las preguntas deben ser presentadas. Las preguntas son ya muy generales o muy específicas, y el respondente es colocado en una situación en la cual se expresa con grados de libertad relativamente amplios.

Las ventajas que ofrece el abordaje más preciso en un respondente en particular contienen las desventajas de este tipo de entrevista en diseños de investigación, en los que existe la necesidad de entrevistar a una gran cantidad de respondentes. La primera desventaja estriba en la limitación de la comparabilidad entre una entrevista y otra. Hyman (*op. cit.*) señala que en el campo de la antropología, donde esta técnica ha sido utilizada con profusión, diferentes observadores han extraído conclusiones distintas respecto a un mismo fenómeno. Las entrevistas no estandarizadas se corresponden casi exactamente con lo que se llama observación participante, aunque por ésta entendemos el proceso total que puede incluir otras técnicas observacionales en la recolección de los datos.

Las *entrevistas semiestandarizadas* son algo menos flexibles que las no estandarizadas. Aquí existe margen para la reformulación y la profundización en algunas áreas, combinando algunas preguntas de alternativas abiertas con preguntas de alternativas cerradas de respuesta. Por lo general existe una pauta de guía de la entrevista, en donde se respeta el orden y fraseo de las preguntas. Preguntas como ¿qué es lo que más le agradaría que se reforme en su sindicato? que están planteadas para permitir un margen de variabilidad amplio en las respuestas de los sujetos, donde se da a éstos oportunidades para que contesten según su propio margen de referencia, su propia terminología, etc., anotando, por lo general, textualmente las respuestas, las que seguramente serán sometidas *a posteriori* a un análisis de contenido.

Las *entrevistas estandarizadas* y los cuestionarios son prácticamente la misma cosa, solamente que se habla de entrevista estandarizada en situaciones en las que el cuestionario se aplica por un entrevistador que leerá las preguntas a un respondente. Las preguntas son presentadas exactamente como figuran en el cuestionario y en su mismo orden. Las preguntas han sido determinadas por el investigador, no permitiéndose por lo general que el entrevistador refrasee o introduzca modificaciones. Las preguntas pueden ser, y por lo general están, “cerradas”; esto es, se le proporcionan al sujeto alternativas de respuesta donde debe seleccionar unas u otras, ordenarlas, expresar su grado de acuerdo o desacuerdo, etc. La entrevista estandarizada ofrece algunas ventajas sobre los otros tipos de entrevista ya que: *a)* permite hacer comparable la información proveniente de distintos sujetos; *b)* facilita la medición, que varía en función directa al grado de estandarización de la pregunta; *c)* aparece como más confiable en la medida en que existe una constancia en los estímulos; *d)* minimiza los errores que se puedan introducir en el refraseo de preguntas; *e)* finalmente, en términos de costos de tiempo y facilidad de procesamiento de los datos e interpretación, la entrevista estandarizada es insuperable.

A estas ventajas le corresponden también desventajas: *a)* existe el problema semántico; por más estandarizada que esté una pregunta, no es posible siempre estandarizar el significado que tiene cada pregunta para distintos respondentes; *b)* una desventaja adicional está dada por la poca flexibilidad, lo que puede llegar a comprometer la situación de comunicación.

En los párrafos subsiguientes y en un capítulo especial, nos concentraremos con mayor detalle en los cuestionarios.

Los *cuestionarios* son pues similares a las entrevistas estructuradas. El instrumento de recolección de datos, que es el cuestionario, será sometido a un proceso de control denominado prueba previa (o *pretest*). De ese proceso de control resultará un cuestionario final, que será utilizado en el paso siguiente, o sea, el de la recolección de datos. Podemos hablar entonces de un *cuestionario inicial* y de un *cuestionario final*.

Para efectuar la prueba previa se necesitará entrevistadores y una muestra extraída del universo que se trata de estudiar. Estas dos cosas serán analizadas más adelante en el texto.

Cuestionario inicial

Desarrollamos aquí principalmente los tipos de cuestionarios en los que se utilizan encuestadores para obtener los datos. Existe otro tipo de cuestionario autoadministrado en el que los propios sujetos leen y registran sus respuestas.

El cuestionario está compuesto de *preguntas, espacio para registrar las respuestas y espacio para registrar la entrevista como unidad*.

Denominamos *espacio para registrar la entrevista como unidad* aquella parte del cuestionario que contiene los datos para ubicación de los sujetos, el número que recibirá para su identificación, los datos referentes al encuestador, fecha en que ha sido realizada la entrevista, su tiempo de duración, etcétera.

Por *espacio para registrar las respuestas* entendemos una distribución particular en la superficie de la página. A veces es conveniente (para facilitar tanto la lectura y escritura por parte del entrevistador como la tarea de análisis de las respuestas) disponer, por ejemplo, las preguntas a la izquierda de la página y el espacio para las respuestas y comentarios del entrevistador, a la derecha. Las preguntas pueden ser de diferentes tipos. Hay preguntas abiertas y cerradas; a veces se aplican también escalas como la Guttman o la Likert. Una pregunta en el cuestionario puede corresponder a una variable, aunque puede también pertenecer a una de las dimensiones de la variable o a un indicador. Para facilitar las respuestas de los *R*, muchas veces es conveniente reunir las preguntas sobre una misma área de interés, en “baterías” de preguntas (se puede consultar más información sobre estas áreas en los capítulos sobre el cuestionario y sobre escalas).

Veamos ahora de dónde salen las preguntas que el investigador va a incluir en su cuestionario. Lo más inmediato, y no tan obvio, es que aquéllas deben hacer referencia a lo que el investigador quiere estudiar. Las fuentes para la obtención son, entre otras, las constituidas por el propio investigador y que se deriven de su sistema conceptual. En

segundo lugar, los cuestionarios realizados por otros investigadores. En las entrevistas cualitativas encontrará también suficiente material. Lo mismo ocurre con la consulta de la bibliografía y los documentos. Por último, haciendo un análisis de contenido del material de discursos, de contenido de artículos de prensa, artículos de revistas, diarios, comentarios de radio y televisión, que estén relacionados con la temática de la investigación, es posible obtener buenas preguntas.

El problema fundamental de la construcción de preguntas es si la pregunta realmente representa la variable, es decir, si el contenido puede servir como indicador de la variable. En el caso de una escala, queremos también saber si el contenido de ella indica la variable. Además, queremos saber si dentro de la escala todos los ítems indican la misma cosa (cualquiera que sea). Se puede, en el proceso del *pretest*, verificar esto comparando las preguntas, estudiando las distribuciones de respuestas y haciendo un análisis de consistencia interna de los ítems en la escala (ver la escala Likert).

Muestra para el pretest

Cuando el investigador tiene el cuestionario inicial listo con todos sus detalles, procede con el *pretest*, es decir, a aplicar ese cuestionario en una pequeña muestra de la población que va a estudiar. Esta muestra para el *pretest* es diferente de la muestra que vamos a aplicar en el paso de recolección de datos, en el sentido de que ahora se trata de obtener una muestra pequeña del medio que estamos estudiando y que no coincida con la muestra “final”, en el sentido en que una persona que aparezca en el *pretest* no lo haga nuevamente en la muestra final.

En el caso en que nuestro campo de estudio sean todos los trabajadores de una fábrica se puede escoger la muestra para el *pretest* de los trabajadores de una fábrica similar en lo que se refiere a alguna variable fundamental. Lo importante en esto es que la muestra del *pretest* se relacione lo más posible con la muestra final y que el número que se tome incluya a unos 30 individuos. Con este número será posible hacer un análisis cuantitativo de los datos obtenidos en el *pretest* al lado de una inspección cualitativa —es decir que se pueden estudiar las distribuciones de las categorías de respuestas, calcular promedios y estudiar dispersiones y también hacer simples cruces bivariados y además ver si los ítems funcionan o discriminan—. Todo esto no será posible si aplicamos el *pretest* a unas 5 u 8 personas, por ejemplo.

Si el investigador dispone eventualmente de recursos suficientes, no perjudica en nada el análisis si el *pretest* es ampliado dos o tres veces; por el contrario, esto acrecienta la exactitud del procedimiento. Por supuesto, el tamaño del *pretest* está ligeramente ligado al tamaño del estudio final. Si el estudio final está calculado para 300 individuos el *pretest* es demasiado largo si es de 150 individuos, etcétera.

Entrevistadores para el pretest

Es recomendable que el investigador mismo efectúe una gran parte de las entrevistas del *pretest*. Si después va a emplear un gran número de entrevistadores, es recomendable

también que algunos investigadores asistentes y los mejores entrevistadores participen en el trabajo de campo del *pretest*. Las razones para esto son las siguientes: hay todavía tiempo para aprender algo nuevo con respecto al problema en estudio, que en muchas ocasiones puede llevar a la reformulación de la hipótesis; pueden introducirse nuevos indicadores; algunas preguntas pueden tener necesidad de ser reformuladas, etc. En el *pretest* el cuestionario inicial no está todavía tan “cerrado” como el cuestionario final; algunas preguntas están abiertas, las que luego del análisis del *pretest* se presentarán como preguntas con alternativas cerradas en la etapa de recolección de datos; en general, toda la entrevista requiere mucha más sensibilidad durante el proceso que cuando se llega al cuestionario final. Por ello, sólo los entrevistadores con bastante experiencia pueden participar juntos con el equipo de investigadores.

Organización y evaluación del pretest

Si unas 5 personas efectúan de 6 a 10 entrevistas, la operación puede estar terminada en un plazo de dos días. La evaluación del *pretest* se hace inmediatamente después que se han efectuado las entrevistas. Algunos puntos que deben profundizarse son los siguientes:

Examen del cuestionario como un todo y en cada una de sus partes:

- ¿qué reacción ha tenido el entrevistado con respecto a la entrevista?
- ¿en qué forma puede lograrse una mayor motivación de parte del *R*?
- ¿cuál es la hora (o día) más oportuna para llevar a cabo las entrevistas, y dónde? (en el lugar de trabajo, en la casa, etc., dependiendo del tipo de entrevista).
- ¿cuál debe ser la longitud de la entrevista?
- ¿en qué orden deben colocarse los diferentes grupos de preguntas (baterías)? Por ejemplo: preguntas sobre deportes, preguntas personales, preguntas sobre aspectos económicos o políticos.
- ¿existe alguna posibilidad de que el orden de los grupos de preguntas dañe el resultado de la entrevista? (Por ejemplo, colocar una pregunta sobre ingreso al comienzo del cuestionario. Ver el capítulo sobre el cuestionario.)

Examen de cada una de las preguntas:

- ¿es necesaria la pregunta?
- ¿podría reformularse?
- ¿podría agregarse una pregunta suplementaria?
- las respuestas alternativas que se han dado son: ¿suficientes? — ¿demasiadas? — ¿no se aplican?
- ¿hay suficiente espacio? ¿Están escritas de modo que se eviten confusiones?

Al mismo tiempo, o más tarde, hay que estudiar la frecuencia de la distribución de las

respuestas (distribuciones marginales). En relación con la forma de la distribución, se deciden las respuestas alternativas finales. A través de este estudio podemos decidir cuándo es necesario presentar la pregunta en una forma dicotomizada o tricotomizada o si dejaremos más posibilidades de respuesta o la pregunta directamente en forma “abierta”. Si se presentan algunas dudas, es preferible presentar una completa escala de alternativas, ya que una dicotomización o tricotomización puede hacerse siempre después, pero lo contrario no es posible; es decir, si sólo dejamos dos alternativas en la escala, no podremos producir después cuatro, en caso de necesitarlas. En relación con esto, se hace una simple tabulación cruzada de las variables más importantes, con el objeto de obtener la primera verificación del resultado de la hipótesis. Las tabulaciones cruzadas están limitadas a cuadros bivariados, debido al número restringido de casos.

A menudo no se presentan a los *R* en el *pretest* las preguntas con distintas categorías de respuestas, sino más bien preguntas abiertas. Esto ocurre cuando el investigador no cree conveniente exhibir las diferentes respuestas entre las cuales el *R* debe elegir, pensando en que de hecho el sujeto o bien no sabe nada acerca del problema, o que indicándole las posibles respuestas podría distorsionar la imagen que sobre el objeto tiene el sujeto. En el primer caso, las respuestas deben ser analizadas después y, si se cree conveniente, clasificar las diferentes respuestas en alternativas que figurarán en el cuestionario final con el objeto de facilitar tanto el trabajo del entrevistador como la codificación posterior.

Sin embargo, muchas veces el investigador entrega ya en el *pretest* las preguntas cerradas, en cuyo caso se instruye a los entrevistadores para que lean las alternativas y soliciten a los *R* que elijan la respuesta que consideren más exacta o que se ajusta más a su manera de pensar.

El último paso en el proceso de evaluación consiste en ordenar y enumerar las preguntas y tipificar el cuestionario nuevamente. Antes de esto, el equipo tiene que decidir si es necesario efectuar un segundo *pretest* o si con el primero se ha llegado a resultados satisfactorios. Ahora vamos a dejar la parte del *pretest* y a pensar más en otros requisitos que deben prepararse con la finalidad de proceder a la recolección final de datos (paso III).

Necesitamos una *muestra* para el trabajo de campo que no tiene que estar mezclada con la muestra del *pretest*; necesitamos también *entrenar a los entrevistadores* frente al trabajo de campo. Este último trabajo se hace en relación con el cuestionario final que se debe tener listo después de la evaluación del *pretest*.

Diseño de la muestra

Si fuera posible, debiera investigarse toda la población. En la mayoría de los casos, la población es muy numerosa y el presupuesto del investigador, limitado. Por lo tanto, debe efectuarse una representación al azar de la población. El tipo de diseño de la muestra que debe utilizarse depende del objetivo del estudio, y de la unidad de análisis que se utilizará, de su distribución geográfica y de la distribución de las otras variables

fundamentales. En uno de los capítulos de este libro se entrega una corta introducción referente a los principales tipos de muestras probabilísticas y no probabilísticas.

Hay varias maneras de construir muestras, y la pregunta que tiene que responder el investigador es qué representa realmente la muestra para su estudio. En caso de que el estudio sea descriptivo y queramos predecir la distribución en la población (valores paramétricos) la muestra debe ser probabilística. Pero cuando el estudio tiene como objetivo establecer las relaciones entre variables de tipo exploratorio, la probabilidad no es el único e indispensable requisito.

Para mostrar la relación entre la fuerza de gravedad, tiempo y longitud, Newton no necesitó recurrir a una muestra representativa de todos los objetos pesados del mundo, para incluirlos en una situación experimental. Muchas veces da más seguridad de predicción el hecho de mostrar la misma tendencia en una muestra pequeña una y otra vez durante distintos periodos que elaborar una investigación grande que explique todo.

No es necesario que el investigador domine, en un estudio sociológico, cada aspecto teórico de la muestra, pero sí es un requisito indispensable la cooperación de un estadista en el caso en que sea necesario. Desgraciadamente los diseños de muestras demasiado complejos no pueden ser aplicados, puesto que las informaciones fundamentales en las cuales la muestra ha sido tomada son incompletas, impracticables, etcétera.

Entrenamiento de los entrevistadores

En los países en que los entrevistadores para hacer estudios no son profesionales, el investigador tiene que seleccionar y entrenar un buen equipo. El proceso de selección y entrenamiento de los entrevistadores puede empezar junto con el *pretest* (en el cual, además, pueden intervenir algunos entrevistadores). Normalmente el investigador encuentra las personas apropiadas entre los estudiantes de su misma y/o de diferentes disciplinas.

Si suponemos que las personas no poseen experiencia anterior en entrevistas, el primer paso es procurarse un cuestionario estándar para el entrenamiento. Este cuestionario incluye diferentes tipos de preguntas, hojas sueltas para responder e instrucciones para los entrevistadores. A éstos deben dárseles algunas lecciones acerca de la teoría de la entrevista y de la aplicación de los cuestionarios. En seguida, hay que reunirlos en pequeños grupos, con el cuestionario de entrenamiento de entrevistas que ha preparado. Luego, hacer que uno entreviste a otro ante el investigador que sigue el procedimiento cuidadosamente (en suma, un juego de *role-playing*).

Hay equipos técnicos más avanzados, que se llaman laboratorios de entrevistas. Estos “laboratorios” pueden consistir en algunas salas especialmente preparadas con micrófonos y amplificadores y también una sala de conferencias, con un espejo *one way*. Naturalmente, esta forma de preparar a los entrevistadores pertenece más a instituciones grandes en universidades. Estos “laboratorios” permiten llegar a cursos regulares en los cuales el participante recibe un diploma y tiene oportunidad de ingresar en un grupo profesional de entrevistadores con su propio código de ética, etc. Sin embargo,

reduzcamos las aspiraciones a la versión más accesible presentada: el investigador efectúa una selección de entrevistadores en esta primera visita de entrenamiento. Ni las personas demasiado tímidas ni las que presumen de su excelencia como entrevistadores llegan a ser buenos entrevistadores. Toda persona que trate de predisponer la entrevista a uno u otro lado no es recomendable. Para predecir si una persona va a ser buen o mal entrevistador, una experiencia vasta como entrevistador y los estudios teóricos de la entrevista por parte del investigador serán de gran ayuda.

Habrà siempre un relativo grado de ventaja en la tarea de llevar a cabo las entrevistas a largo plazo. Entrevistadores que son excelentes en un tipo de estudio pueden ser mediocres en otro; y entrevistadores que han sido mediocres en el comienzo de la colección de datos pueden resultar las mejores personas para convencer a los *R* que en un principio rehusaron ser entrevistados.⁵

Cuando el investigador tiene un equipo de entrevistadores capaces y bien entrenados, el cuestionario de la investigación final se presenta junto con las instrucciones para los entrevistadores. Llevarà algùn tiempo al entrevistador manejar todos los formularios y el cuestionario. Antes de empezar el trabajo en el campo, el equipo de investigación debe cerciorarse de que los entrevistadores han comprendido estas instrucciones y dominan el cuestionario, pues de lo contrario sólo conseguirían verse confusos o inseguros (ver también el capítulo sobre la organización del trabajo de campo).

PASO III: LA RECOLECCIÓN DE DATOS. EL TRABAJO EN EL CAMPO

El trabajo en el campo debe planearse y llevarse a cabo con precisión y seguridad. En un estudio de tamaño normal, digamos de unos 500 casos, la parte principal de las entrevistas (un 70%) debe ser obtenida en un plazo de cinco días después de su comienzo en el campo de trabajo. El número de entrevistadores y recursos deben calcularse sobre estas cifras. Por supuesto que depende mucho del tipo de estudio que se lleve a efecto. Algunas entrevistas tienen lugar en áreas rurales, otras en parte de ciudades que no tienen muchos medios de comunicación. De todos modos, no es necesario justificar una demora demasiado prolongada de la recolección de datos, por causa de dificultades que pudieran producirse. En muchos estudios es fácil encontrar dificultades de diferentes tipos.

Si el estudio se va a llevar a cabo en un pueblo, por ejemplo, se puede enviar allí a unos 30 entrevistadores por un periodo de cinco días, que vivan cerca de la central de entrevistas (podría, perfectamente, ser un hotel), con instalaciones de teléfonos y comunicaciones y en un lugar céntrico. *Los entrevistadores están precisamente para cierto propósito y para eso se les está pagando.*

Los *R* se dividen de tal modo que todos queden próximos. A cada sector geográfico se le da un entrevistador, etc. Los entrevistadores, por ejemplo, tendríamos que informar por cada cinco entrevistas a la central, con objeto de que dichas entrevistas sean debidamente verificadas, etc. Serán cinco días sobrecargados, pero posiblemente el estudio será realizado de una manera eficiente.

Un problema en toda recolección de datos primarios es el tipo de dificultades que se presentan para obtener el total de las entrevistas. Hablamos de pérdida de datos. La pérdida de datos puede ser en la forma de pérdida de la entrevista entera o de parte de ella. En el primer caso, es la pérdida de la unidad, y en el segundo, es un vacío en la matriz de datos. Nos preocuparemos aquí de la primera posibilidad.

La gente puede rehusar las entrevistas por diversas causas imaginables. En una encuesta al azar de 500 *R* una serie de inconvenientes puede aparecer. No es fácil tomar simplemente a alguna otra persona para sustituir los individuos que rehúsan la entrevista. Sin embargo, existe cierto porcentaje de entrevistas que el investigador obtendrá fácilmente, digamos alrededor del 70-75%, y hay un 20% difícil de obtener. Mientras más alto sea el porcentaje que deseemos obtener, más alto será el costo y el trabajo que tendremos en las últimas entrevistas. Teóricamente, podríamos viajar a Europa, por ejemplo, para obtener las entrevistas de algún *R* que está allí de visita durante el tiempo de nuestra investigación. El costo, sin embargo, y el tiempo ponen límite a los altos porcentajes y tenemos que continuar el estudio sin lograr el 100%. Todo lo que se obtenga superior al 90% es aceptable, si usamos los estrictos criterios mencionados.

Es necesario recordar algo que justifique el rigor de nuestro criterio: el último grupo de 10% que obtenemos es un grupo socialmente especial en el sentido de que son personas difíciles de ubicar. Podría tratarse de una específica forma de personalidad que rehúse todo lo que se haga en sociedad, incluyendo las entrevistas; podrían ser personas inestables y difíciles de localizar; podría tratarse de enfermos, etc. Esto es, probablemente un grupo más interesante de estudiar con respecto a algunas de nuestras variables y por eso no puede ser sustituido por otro tipo de personas.

PASO IV: EL PROCESAMIENTO DE DATOS

Bajo el título de procesamiento de datos encontramos subtítulos como “codificación”, “perforación”, “reproducción”, etc. En la siguiente representación introductora del proceso de un estudio tipo *survey*, tocaremos muy brevemente algunos aspectos y dejaremos que el lector investigue sobre el particular en otros capítulos del libro.

Codificación

Codificación es el traslado de categorías de respuestas a un lenguaje simplificado (cifras) con el objeto de efectuar el proceso de análisis dentro del espacio adecuado en la tarjeta IBM u otras tarjetas perforadas.

En el proceso de codificación usamos una lista de códigos que, como un diccionario, da la correspondiente connotación de una categoría en otro idioma. Damos a continuación un ejemplo de código:

<i>Pregunta núm.</i>	<i>Variable</i>	<i>Columna núm.</i>	<i>Código</i>	<i>Alternativas</i>
56	Ingreso familiar mensual en escudos (Chile), 1968	71	1	Menos de E° 400
			2	> de E° 400 < E° 800
			3	> de E° 800 < E° 1.200
			4	> de E° 1.200 < E° 1.600
			5	> de E° 1.600 < E° 2.000
			6	> de E° 2.000 < E° 3.000
			7	> de E° 3.000

Se puede estudiar más en detalle la codificación en el capítulo que trata el asunto.

Perforación de las tarjetas

Cuando los datos crudos han sido codificados en una hoja especialmente dibujada para estos fines o simplemente se han codificado en el cuestionario en uno de los márgenes, el próximo paso es perforar las cifras en las columnas de la tarjeta de datos. De este modo, obtendremos para cada individuo una tarjeta (o un grupo de tarjetas si el número de variables es numeroso); dentro de la tarjeta, cada columna o grupo de columnas representará una variable o una “pregunta” en el cuestionario; en cada columna la posición corresponderá a la respuesta alternativa escogida por el R.

La perforación es un procedimiento relativamente rápido. Las tarjetas necesitan también ser verificadas y pueden ser reproducidas con el objeto de obtener una o más copias del juego inicial de tarjetas. Todo este trabajo es efectuado por máquinas especiales que pueden ser alquiladas comercialmente. En muchos departamentos universitarios y otras instituciones científicas existe, además, una oficina especial para computaciones que ofrece esta clase de servicios.

Una pregunta es de especial importancia para el investigador con respecto al procesamiento de datos: ¿cuándo es recomendable usar computadoras, cuándo una clasificadora y cuándo simples técnicas, como la técnica con tarjetas Mc Bee, o simplemente hacer el análisis directamente en tablas, sin transferir el material a tarjetas? Esto depende del número de casos que tenga el estudio, el número de variables utilizadas y de la complejidad del análisis que se quiere utilizar.

Para un *pretest* con más o menos 30 a 50 casos, las tarjetas Mc Bee habitualmente ofrecen la mejor solución. Para un estudio de alrededor de 200 a 600 casos con un número de variables no muy grande y con el fin de no complicar el procedimiento analítico, las tabulaciones cruzadas en un simple distribuidor de tarjetas (clasificadora) es el más económico de los métodos en tiempo y dinero y permite cambios inmediatos en el plan de análisis. Para materiales más amplios, una computadora ofrece la mejor alternativa.

Usando máquinas clasificadoras corrientes, máquinas estadísticas o servicios de computadoras, el investigador obtiene las distribuciones de sus variables con sus

respectivas características, como promedios, dispersiones, etc. A partir de esta información, procede a diseñar un plan para el análisis.

PASO V: ANÁLISIS

A través de la información obtenida de los datos, el investigador continúa su labor haciendo un proyecto de tabulaciones cruzadas, incluyendo correlaciones, tests de significación estadística y otros procedimientos analíticos que desee aplicar y que se corresponden con su sistema de hipótesis.

PASO VI: PRESENTACIÓN

La presentación final de los hallazgos deberá ser efectuada de tal manera que muestre sus resultados en forma sistemática y compacta. Asimismo, la presentación debe incluir un apéndice metodológico en el cual el cuestionario y otros detalles en la recolección de datos y proceso se encuentren debidamente explicitados. Esto se realiza con objeto de hacer más fácil a otros científicos la investigación del sistema empleado, efectuar réplicas y reconstrucciones de preguntas y cuestionarios y, de esta manera, avanzar en la teoría.

COMENTARIOS FINALES

Si vemos el diagrama al inicio de este capítulo, en el que se aprecian los diferentes pasos que sigue una investigación, notamos una flecha entre el paso V y el paso I. Esta flecha representa la interconexión del paso de análisis con el paso lógico de la formulación de nuestra hipótesis. Este paso representa la interpretación y consiste en confrontar los resultados hipotéticos que se derivan de la teoría con los hallazgos empíricos. La confirmación de una hipótesis o la refutación de la misma tienen consecuencias importantes; y a este respecto es importante justificar, por ejemplo, tanto la presencia como la ausencia de correlaciones. Es señal de poca ética profesional presentar en el informe de investigación solamente aquellos hallazgos que coinciden con nuestras hipótesis o, peor, con nuestros puntos de vista, e ignorar los que no coincidan.

Cuando analicemos en el próximo capítulo el proceso de investigación, notaremos la manera de ver una investigación tipo *survey* como un pasaje entre el sistema de hipótesis y el análisis con el trabajo de corrección de datos ubicado en una forma intermedia. Es decir, al suprimir el paso I y el V, no habrá una investigación explicativa. Suprimiendo los pasos intermedios II, III y IV, no habrá verificación empírica del I y no se puede tampoco efectuar el paso V.

II. EL PROCESO DE INVESTIGACIÓN

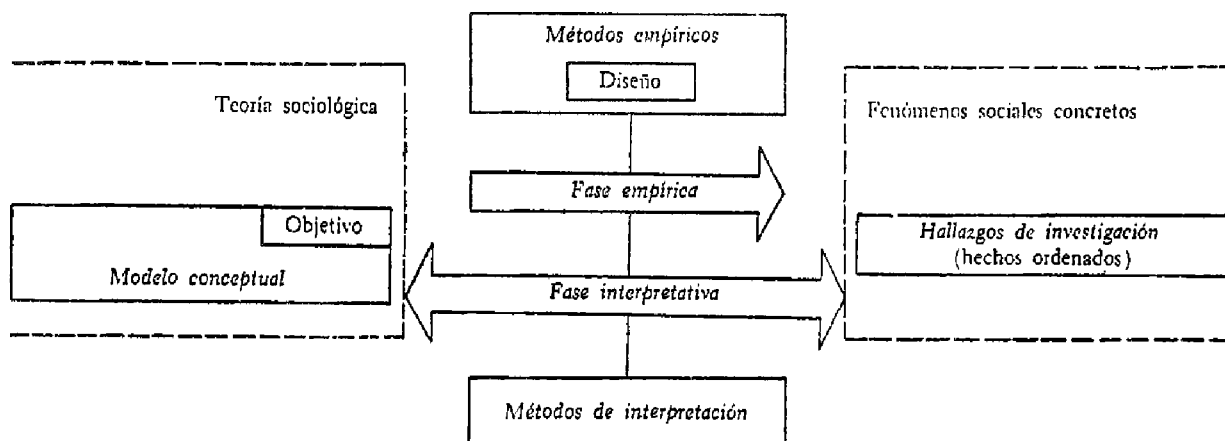
JORGE PADUA

EL PROCESO de investigación está constituido por una serie de partes íntimamente relacionadas. Del conocimiento de tal interconexión y de su manejo correcto a lo largo de toda la dinámica del proceso dependerá el resultado sustantivo de la investigación misma.

Las partes principales del proceso aparecen enumeradas en el diagrama del capítulo anterior. De esas partes, son dos las que desempeñan un papel fundamental: el marco teórico y el análisis. La recolección de datos en sí, cualquiera que sea su grado de sofisticación, es un instrumento que garantiza el paso del marco teórico a la verificación del mismo.

La afirmación expresada en el párrafo anterior se hace más explícita cuando se entiende que el marco teórico *inicia* el proceso y da lugar a una problemática expresada *a priori*, en la forma de un conjunto de proposiciones que, si se presentan en forma aislada, caracterizarán un estudio como descriptivo; si están interconectadas o interrelacionadas con otras proposiciones, lo caracterizarán como explicativo.¹

Todo problema de investigación comienza siempre como un problema de teorización.² La primera tarea del investigador es la de la codificación de la realidad, y esa codificación sólo puede ser realizada según una teoría (implícita o explícita). Matilda Riley³ presenta un diagrama del proceso de investigación que concentra el interés en las tres partes principales del proceso: la teoría, los métodos empíricos para la recolección y la realidad o fenómenos sociales concretos. El diagrama aparece a continuación.



La *teoría sociológica* es el conjunto de proposiciones y definiciones extraídas de la

realidad social y que explican los fenómenos sociales concretos. El *modelo conceptual* es construido a partir de términos generales, definiciones y supuestos de la teoría o de una porción de ella. Este modelo señala cuáles son los problemas más significativos, las maneras como se seleccionarán los datos, la selección del diseño más correcto, así como la búsqueda de orden o de patrones entre ellos y la interpretación de los hallazgos de investigación.

El *objetivo de la investigación* señala los elementos en el modelo que van a ser investigados; pero, repetimos, es el modelo el que señala los supuestos que están por detrás de los elementos. Sin embargo, los objetivos de la investigación pueden ser también la selección de un modelo.

El *diseño de investigación* se refiere al conjunto particular de métodos seleccionados por el investigador tanto para la búsqueda de nuevos hechos como para la determinación de sus conexiones. Aquí se decide cómo se van a seleccionar los datos, cuáles serán los métodos analíticos, cómo se va a formular el problema, qué tipos de instrumentos específicos se van a utilizar, cómo se va a realizar el *pretest*, etcétera.

En la *fase empírica* el investigador es guiado por la teoría sociológica, y el modelo hacia los fenómenos sociales concretos que, en términos de hechos sociales, contrastarán sus hipótesis teóricas.

En la *fase interpretativa* se comparan los hechos con su teoría inicial examinando las consecuencias que tienen para la teoría la comprobación o refutación de las hipótesis.

Así, pues, el problema básico es un problema de teorización. Sin embargo, cuando se habla de teorización, por lo general se hace referencia a dos niveles distintos, pero relacionados: a) el *nivel conceptual* y racional de conceptualización, que define el significado nominal de una ciencia, y b) el *nivel de teorización* para la investigación, es decir, el nivel empírico, que le da el significado operacional a una ciencia.

Ambos niveles de teorización están relacionados en el sentido de que debe existir una correspondencia real entre el significado nominal de una ciencia y su significado operacional; es decir, debe existir un isomorfismo entre ambos niveles para que el resultado de la operación científica sea conocimiento.

Hay todo un problema epistemológico en las ciencias sociales que proviene de la falta de correspondencia entre el significado nominal y el significado operacional, dificultad a la que se han propuesto algunas soluciones parciales, que desafortunadamente por su carácter de parciales no han resuelto el problema de una manera satisfactoria. Por el carácter práctico e introductorio del manual, tocaremos muy de lado la consideración de las soluciones propuestas, prefiriendo enfrentar el problema desde una perspectiva diferente.

Los problemas de teorización pueden ser vistos desde la teoría hacia la investigación, o desde la investigación hacia la teoría. Si bien esto no significa desplazar el problema de la ausencia de isomorfismo entre los niveles de significación y los niveles operacionales, sí implica un desplazamiento de la importancia relativa que se da a los enfoques en la resolución del problema. Si uno mira el problema desde la teoría hacia la investigación, la preocupación deviene casi exclusivamente epistemológica; si la

perspectiva va de la investigación a la teoría el problema es más técnico, más pragmático. No se trata de dar a una u otra perspectiva un *status* superior o inferior. Si tomamos la perspectiva técnica en este texto, lo hacemos conscientes de que, si bien las soluciones se concentrarán en problemas operacionales, éstas dependerán —sin ninguna duda— de los resultados y del nivel alcanzado por la disciplina en la resolución de sus problemas a nivel sustantivo.⁴

Por operacional debe entenderse una solución que corresponda al nivel de las necesidades pragmáticas de realizar una investigación empírica, sin dejar de lado los problemas de significado nominal. La aclaración es válida, sobre todo para diferenciar operacional de operacionalismo; esto es, de la práctica que consiste en suprimir los significados nominales en beneficio de operaciones específicas y uniformes.

DISTINTOS TIPOS DE INVESTIGACIONES

La investigación científica tiene como sus objetivos teóricos más generales dar respuestas inteligibles, confiables y válidas a preguntas específicas o problemas de investigación. Las respuestas se dan por lo general en términos de qué (o cómo), dónde, cuándo, de dónde y por qué. Sin embargo, no toda investigación tiene como propósito responder a *todos* los interrogantes, existiendo la posibilidad de que se trate de responder solamente a alguno de ellos. Toda investigación comienza, pues, con algún tipo de interrogante, que tratará de ser resuelto en los términos citados.

La *formulación del problema* de investigación es uno de los pasos principales y más difíciles de resolver en cualquier diseño de investigación. El tipo particular de estilo cognoscitivo de una investigación de carácter científico exige del investigador no solamente *claridad* en la formulación del problema a investigar, sino también *especificidad* en términos del tipo de respuesta que se busca a tipos específicos de preguntas.

Si en un comienzo los intereses del investigador pueden ser de carácter muy amplio, en términos de la investigación concreta siempre hay que tener bien claro qué es lo que se está buscando, y el tipo de información que dará la respuesta a sus preguntas. Por ejemplo, el investigador puede tener como área de interés general aspectos vinculados al sindicalismo, e imponerse como objetivos buscar la resolución de algunos interrogantes sobre sus características, organización, etc. Pero para los propósitos de la investigación, el problema tiene que ser formulado de manera más específica, buscando las respuestas al nivel de los cinco interrogantes formulados más arriba. Hay que plantear preguntas tales como: ¿cuántos sindicatos hay?, ¿quién los dirige?, ¿cómo están organizados?, ¿con qué tipo de ideología operan?, ¿a qué tipo de sindicalismo representan?, ¿cuáles son las relaciones con las bases?, ¿por qué se dan esos tipos y no otros?, etcétera.

La bibliografía especializada acostumbra diferenciar los estudios o diseños de investigación —según el tipo de pregunta que el investigador plantee— en estudios *exploratorios*, *descriptivos* y *explicativos*.

Los *estudios exploratorios* son preponderantes en áreas o disciplinas en donde las

problemáticas no están suficientemente desarrolladas, de manera que el investigador tiene como propósito “ganar familiaridad” con la situación antes de formular su problema de manera específica.

Por ello, antes que concentrarse en observaciones particulares, la estrategia de investigación consiste en buscar una mayor dispersión posible en las observaciones. En las ciencias sociales, donde las teorías no están formuladas en forma precisa, los estudios exploratorios son necesarios ya sea para la precisión o examen en profundidad de algunos de los supuestos de la teoría, para la construcción de esquemas clasificatorios provisionales, para detectar algún modelo aún no formulado en forma explícita, o bien para facilitar la generación de algunas hipótesis que serán puestas a prueba posteriormente con algún diseño explicativo. Aquí es conveniente repetir que el investigador siempre opera con alguna clase de “modelo” conceptual, con alguna clase de “teoría” sobre la naturaleza del fenómeno o de la situación; por eso la dificultad principal de la aproximación exploratoria puede residir en la creencia, por parte del investigador, de que puede desprenderse a voluntad de sus estereotipos o preconcepciones.

En la medida en que el tipo de estudios exploratorios involucra por lo general un tipo de “inmersión total” del investigador en la situación, éstos han sido tradicionalmente preferidos por antropólogos y algunos psicólogos sociales, utilizando la observación participante como técnica para la recolección de datos.

Un obstáculo importante para este tipo de estudios puede sintetizarse en el efecto que dos vertientes tienen sobre la validez interna y externa de los resultados de la investigación: a) *problemas de error*, particularmente cuando se utiliza como técnica de recolección de datos algún tipo de observación participante (por ejemplo, se producen modificaciones en la estructura de los grupos, en el comportamiento de los miembros, existiendo escasez de control sobre los sesgos potenciales que el investigador pudiese introducir en el registro de la información, etc.), y b) *problemas de comparabilidad* con otras investigaciones del mismo carácter, y que en general se derivan de la falta de sistematización.

Los *estudios descriptivos* son más específicos y organizados que los estudios exploratorios, ya que las preguntas aparecen guiadas por taxonomías, esquemas descriptivos o tipologías; en estos estudios el interés está enfocado en las propiedades del objeto o de la situación a ser clasificadas en el interior de estos esquemas. Los estudios descriptivos dan por resultado un diagnóstico.

Hans Zetterberg⁵ sistematiza en forma precisa las aproximaciones que enfatizan definiciones y las que enfatizan proposiciones, sistematización que permite una distinción bastante clara en los diferentes significados de la palabra “teoría”, según ésta responda al interrogante de los porqués o al de los cómo; esto es, según enfoque su interés en la explicación o en la descripción.

La palabra “teoría” se reserva para los sistemas proposicionales, es decir, aquéllos donde la unidad es la proposición, y donde las proposiciones ordenadas constituyen el sistema de organización conceptual que, aplicado a las unidades de interés de la disciplina, da por resultado una explicación.

La palabra “taxonomía” se aplica a la interrelación de definiciones; una taxonomía es un esquema de definiciones ordenadas que definen tanto el objeto de interés de una disciplina como las propiedades del objeto o de la situación a los que hay que prestar atención. La aplicación de la taxonomía a un nuevo objeto permite un diagnóstico o una descripción. En un apéndice adjunto a este capítulo presentamos ejemplos de taxonomía en botánica y psicología y una tipología sociológica.

Los *estudios explicativos* dan respuestas a los porqués. La respuesta se ubica dentro de la lógica de las explicaciones científicas, a través de teorías organizadas de manera tal que dan cumplimiento a la condición doble de verificación lógica y de verificación empírica. Lógica en el sentido de consistencia entre las proposiciones que integran el sistema; empírica en el sentido de correspondencia del sistema proposicional con la realidad empírica. En la fase explicativa de la investigación, el investigador separa — como dice Piaget—⁶ lo verificable de lo reflexivo o intuitivo, elabora métodos adaptados a la problemática, métodos que son a la vez de análisis y de verificación.

El concepto de “ley” pasa a desempeñar un papel importante en el aparato conceptual. Mario Bunge⁷ especifica cuatro significados distintos de la palabra “ley científica” que resultan de extrema utilidad tanto para la especificación de los problemas como para la determinación del nivel de preguntas que plantea una investigación de carácter explicativo. La especificación de Bunge sigue de cerca las preguntas que Zetterberg⁸ sugiere que se realicen a una disciplina, para la determinación de su grado de científicidad.

Para la ejemplificación de los cuatro significados de la ley científica, seguimos de cerca a Bunge y su ejemplo de la ley de Newton del movimiento:

1) El primer significado de ley científica (significado ontológico), para Bunge, es la ley como denotando toda *relación constante y objetiva en la realidad*; es decir, la ley como estructura nómica, de pautas invariables al nivel óntico, como pauta objetiva. La ley del movimiento de Newton sería la expresión de una pauta objetiva del movimiento mecánico, pauta que no sería ni verdadera ni falsa sino que simplemente es.

2) El segundo significado de la palabra ley es la ley como *enunciado nomológico* o como enunciado de ley; es decir, una hipótesis general que tiene como referencia la ley entendida en su primer significado. En este caso la ley es una pieza de conocimiento, expresado en forma proporcional. Este tipo de enunciado nomológico nace del hecho de que, en la medida en que no es posible captar las leyes en su significado óntico en toda su pureza, se construyen estos enunciados nomológicos que pertenecen a un modelo de la realidad, modelos que son ideales y que son reconstrucciones cambiantes de las leyes objetivas en el pensamiento científico. Estos enunciados nomológicos son lógicamente necesarios en la medida en que entran en conexión con otros enunciados nomológicos; pero son fácticamente contingentes en la medida en que los enunciados son perfectibles, o solamente válidos en cierto dominio. En el ejemplo de la ley del movimiento de Newton, el enunciado de ley se expresa en la designación de la fórmula “fuerza = masa × aceleración”.

3) El tercer significado de ley es la ley como enunciado *nomopragmático*, que

designa toda regla mediante la cual puede regularse una conducta. En este caso la ley opera como regla de acción, mediante la cual, siguiendo el ejemplo de Newton, es posible predecir o controlar la trayectoria de cuerpos. Aquí la ley denota relaciones invariantes a nivel pragmático, que operan como guías para la acción con fundamento científico, ya que se derivan de los enunciados nomológicos. Esta distinción entre enunciados nomopragmáticos y nomológicos es importante según Bunge, ya que nos permite establecer una distinción entre explicación y predicción. Los enunciados nomológicos constituyen sistemas lógicamente organizados u organizables, que sirven para la explicación. Los enunciados nomopragmáticos, que se derivan de los enunciados nomológicos, son las herramientas para la predicción.

4) Finalmente, el cuarto significado de la palabra ley es la ley como *enunciado metanomológico*, que designa todo principio general acerca de la forma y del alcance de los enunciados de ley. Aquí estamos a nivel de la ley como prescripción metodológica, en la medida en que los enunciados metanomológicos son enunciados explícitos de leyes acerca de las leyes que sirven de guía para la construcción de teorías. El significado de ley está entendido como cierta pauta de la conducta humana en relación con cierto tipo de objetivos y cierto tipo de datos empíricos.

TEORIZACIÓN EN LA INVESTIGACIÓN

Puede hablarse entonces de dos niveles de teorización, el nivel taxonómico (donde se hace hincapié en las definiciones y cuyo resultado será una taxonomía o una tipología o un esquema clasificatorio); y de un nivel teórico proposicional, la teoría propiamente tal, donde se incluye tanto el problema de las definiciones conceptuales cuanto la interrelación de proposiciones con el propósito de explicar.⁹

Los conceptos y constructos incluidos en el nivel de la teoría son abstracciones formuladas a partir de generalizaciones de observaciones particulares, que son definidas nominalmente y que proporcionan el nivel de significado o de sentido del sistema. Los constructos y los conceptos tienen significados similares, siendo los constructos conceptos de un nivel más alto de abstracción que tienen como base conceptos de un nivel de abstracción de bajo nivel. Conceptos y constructos son introducidos en la teoría por medio de definiciones o por medio de operaciones. En el caso de ciencias empíricas se trata de desarrollar un sistema conceptual que tanto a nivel de definiciones nominales como de definiciones operacionales permita el contraste de la teoría con la realidad.

Los conceptos entonces describen los fenómenos y en general tienden a ser de dos tipos: *conceptos categóricos* que son complejos y que se miden a nivel de categorías nominales; y las *variables* que representan dimensiones de los fenómenos admitiendo grados de variación que se miden a niveles ordinales, intervalares o racionales.

El concepto es pues un nombre, hay que agregarle una definición. Cuando se ordenan conceptos y definiciones se obtienen esquemas descriptivos que nos sirven para clasificar o diagnosticar la realidad.

Las *definiciones*¹⁰ son una forma de “explicar”, si atendemos al significado

etimológico de la palabra (desplegar, hacer las cosas planas, más claras, más comprensibles). Pero en las “explicaciones”, por definición, solamente existe un intercambio de símbolos. La definición implica pues un problema semántico, y el objetivo principal es la clarificación de significados. Esta clarificación de significados es el primer paso importante en la clarificación de eventos. Además, y ése es el fundamento de todos los sistemas definicionales, facilitan la comunicación a nivel profesional, homogeneizando el lenguaje.

Se acostumbra diferenciar dos clases principales de definiciones, existiendo dentro de cada una de las clases diferentes tipos:

a) Las *definiciones nominales*. Son simples convenciones lingüísticas que no expresan ningún valor de verdad. Son simplemente una indicación sobre cómo utilizar el lenguaje, y cuyos significados son dados en forma relativamente arbitraria. Este tipo de definición opera solamente en el nivel simbólico y lingüístico, y son juzgadas en relación a su utilidad. Esto es, entre dos definiciones, no es el caso discutir sobre cuál es la más verdadera, sino cuál es la más útil. De las definiciones nominales la más utilizada es la *convencional aristotélica*, en la que se definen los fenómenos por medio de dos atributos: un *genus proximum* y una *differentia specifica*; el *genus proximum* es un atributo que el fenómeno o concepto a definir comparte con una clase más amplia de fenómenos, y la *differentia specifica*, un atributo peculiar a la categoría definida. En resumen, una definición nominal introduce una nueva expresión, ya sea a través de una nueva palabra o de un símbolo o de una frase compuesta por medio de una serie de criterios racionales.

b) Las *definiciones reales*. Operan en el nivel simbólico y además en el nivel referencial. Las definiciones reales son en sí hipótesis que expresan un valor de verdad. Una definición real, para ser válida, necesita ser probada empíricamente como tal. Por consiguiente, las definiciones reales se juzgan en función de su valor de verdad y se establecen mediante la investigación empírica. Lo que el investigador busca mediante una definición real no es solamente lo que la palabra significa, sino cuáles son los referentes del concepto y cuáles son sus propiedades observables.

Una definición real, en un sentido más estricto, es un tipo particular de proposición: una proposición universal afirmativa. Por ejemplo, cuando definimos “hombre” como *bípedo implume*, estamos implicando que cualquier cosa que sea bípeda e implume es hombre, y que cualquier hombre debe ser bípedo e implume. Si bien toda definición real es una proposición, no toda proposición es una definición real, en la medida en que estas últimas no solamente tienen consecuencias lógicas, sino además ontológicas; esto es, tienen significado real y significado empírico. La proposición “todos los hombres son mortales” es verdadera como proposición lógica, pero no es una definición real, ya que no todos los mortales son hombres.

Carl Hempel (*op. cit.*) establece una diferenciación importante entre definiciones reales, definiciones analíticas y definiciones empíricas, que es muy útil, ya que permite hacer más flexible el concepto de “definición real” tradicional, demasiado sujeto a las nociones de atributos y de naturaleza-esencia. Para Hempel las definiciones analíticas

requieren como criterio de validación solamente una reflexión sobre los significados de las expresiones que las constituyen, mientras que las definiciones empíricas requieren el análisis empírico. Más adelante presentamos algunos ejemplos de P. Lazarsfeld, que basándose en las conceptualizaciones de Hempel, sugiere un método para la operacionalización de variables, mediante la idea de pasaje de conceptos a indicadores.

Proposiciones e hipótesis

Cuando se conectan conceptos tenemos un juicio teórico; por ejemplo: a más alta la extracción de clase de un alumno, mayor la probabilidad de que éste tenga un rendimiento alto en la escuela primaria. Estas conexiones entre conceptos están justificadas tanto a *nivel teórico*, cuando el investigador provee las razones acerca de por qué los conceptos deben conectarse de esa forma y no de otra (por ejemplo, en el juicio anterior aduciendo el efecto de la socialización diferencial por clase social, las expectativas de los maestros, etc.), cuanto a *nivel operacional*, es decir, a nivel de cómo se conectan los conceptos, especificando los parámetros del juicio (por ejemplo, especificando si la relación es lineal, curvilínea o potencial; si los coeficientes representan la unidad o una fracción de ella; si la relación es positiva o negativa).

Proposiciones, hipótesis y enunciados se utilizan por lo general como sinónimos, aunque en el sentido técnico pueden significar cosas distintas. Por lo general, los investigadores hablan indistintamente de proposición o de hipótesis, mientras que los epistemólogos utilizan con mayor frecuencia las palabras “enunciado”, “proposición” e “hipótesis” con un carácter más preciso.

Por *proposición* vamos a entender “cualquier generalización que puede probarse como consistente o inconsistente con respecto a otras generalizaciones que forman parte del cuerpo organizado de conocimiento; las proposiciones científicas, además, deben ser sometibles, directa o indirectamente, a la verificación empírica”. Una *hipótesis* es un juicio de carácter conjetural. Para W. I. Beveridge,¹¹ la hipótesis “es la técnica mental más importante del investigador y su función principal es sugerir nuevos experimentos o nuevas observaciones”. Cuando se habla de hipótesis, por lo general el investigador se ubica en una lógica de tipo deductivo, donde las hipótesis son algunas implicaciones de la teoría que esperan verificación.

Para Zetterberg (*op. cit.*) las proposiciones teóricas o de alto valor informativo son aquellas que “pueden ser probadas como incorrectas por un gran número de distintas maneras”, esto es, las que dan cuenta de una gran variedad de eventos. Galtung define las hipótesis como conjunto de variables interrelacionadas y una teoría como conjunto de hipótesis interrelacionadas.¹²

Como puede verse, no existe distinción sustancial entre la definición de proposición de Zetterberg y la definición de hipótesis de Galtung. En el fondo están hablando de la misma cosa: Galtung, desde la perspectiva de la matriz de datos, y Zetterberg, desde las teorías de la axiomatización. Para los propósitos de la investigación empírica, los criterios de Galtung son algo más operacionales que los de Zetterberg. El primero

sostiene que una hipótesis es una oración o sentencia acerca de cómo se distribuyen un conjunto de unidades en un espacio de variables. Para el lenguaje de la matriz de datos, las hipótesis tienen la forma siguiente: *a)* se da una unidad (mesa, clase social); *b)* una o más variables (color, participación política); *c)* las variables tienen valores (rojo, azul, etc.; alta, media, baja); *d)* puede hacerse referencia a un conjunto de unidades con más de un elemento, y *e)* puede referirse a un atributo o propiedad de la unidad.

A diferencia de Zetterberg, Galtung concibe hipótesis de una variable (descriptiva): esta mesa es verde y alta; de dos variables: a mayor clase social, mayor participación política; o multivariatas, en cuyo caso comenzamos a hablar de sistema de hipótesis, tal como, por ejemplo, aparece en el diagrama del capítulo I, donde además de hacer referencia a un atributo o propiedad de una unidad, establecemos vínculos —que pueden ser causales— ya sea entre distintas propiedades de una unidad, ya sea de interrelación entre distintas unidades y propiedades.

Hasta aquí lo que importa destacar es la idea de explicitar las proposiciones o hipótesis. Es aconsejable hacer algunos esfuerzos para llegar a un máximo de axiomatización de la teoría, entendida ésta como sistema proposicional o de hipótesis.

La explicitación de las proposiciones y de las definiciones, y su tratamiento teórico a través de la reducción de matrices (combinando proposiciones con definiciones o proposiciones con proposiciones), tienen la ventaja de hacer visibles todas las ideas implícitas en algunas ideas dadas.¹³

Una vez explicitadas las proposiciones, el siguiente paso es la verificación de éstas (si se trata de hipótesis/proposiciones aisladas) o de las interrelaciones entre las proposiciones, en esquemas multivariatos.

La verificación en el caso del conjunto de proposiciones interconectadas consiste en ver cuán bien (cualitativa y/o cuantitativamente) se ajusta esta estructura con la estructura empírica que indican los datos.

Aunque la idea es simple, de ningún modo es sencilla. Llegar a la formalización no es tarea nada fácil, principalmente porque nuestras disciplinas están en etapas embrionarias en lo que a especificación de estructuras causales se refiere (y sobre el problema de la causalidad misma aún no se ha dicho la última palabra).

Por otra parte, no hay todavía acuerdo sobre qué se entiende por algún concepto dado o por algún otro. Más técnicamente, como expresamos líneas arriba, no hay taxonomías completas y es insuficiente el trabajo que se hace por desarrollarlas y, como sabemos, poco puede la teorización sin buenas taxonomías y las taxonomías sin buenas teorizaciones.

Pero es una excelente manera de operar y hay muchos investigadores que no la utilizan. Consideremos nuestro medio: es práctica muy utilizada el que, en base a algunas inquietudes más o menos difusas, y en todo caso bien implícitas, se recogen primero los datos y luego se fabrica la teoría, según lo que den los cuadros. En el mejor de los casos, en estas investigaciones se publicarán algunas caracterizaciones generales. En el peor, las investigaciones terminan en la recolección de datos, acumulando cuestionarios o tarjetas perforadas y guardándolos en un rincón de alguna oficina.¹⁴ Y es

por ello que volvemos a señalarlo: no se insiste como es debido en la teorización y en la dinámica del proceso de la investigación.

CONCEPTOS, INDICADORES, ÍNDICES

Como último punto, como ejemplo de una de las maneras de trabajar, queremos exponer la operacionalización de algunas de las dimensiones de la personalidad autoritaria,¹⁵ a través de las ideas desarrolladas por Lazarsfeld y colaboradores¹⁶ sobre el pasaje de los conceptos a los índices empíricos. Podría decirse que es el paso de un concepto a una variable.

Supóngase un sistema de hipótesis (o de proposiciones) o un número de proposiciones que pretendan describir o explicar aspectos relacionados con elementos de la personalidad que puedan predisponer formas de reacciones hostiles frente a ciertos grupos raciales. A partir de esta idea general, es necesario comenzar con la especificación y organización de la forma como se abordará el tema de investigación.

Lo que aquí nos interesa es que, una vez que uno haya explorado la literatura existente, y se hayan formulado las proposiciones, se debe tomar cada uno de los conceptos que involucran las proposiciones y proceder a operacionalizarlos para poder aprehenderlos en forma empírica. Retomemos el caso de la “personalidad autoritaria”. La idea principal de los autores es que la personalidad humana funciona como un todo, de manera que sus convicciones políticas, económicas y sociales son el reflejo de aquellas tendencias profundas de su personalidad.

La idea de Lazarsfeld es que el paso de los conceptos a los índices incluya cuatro etapas: A) La imagen inicial. B) Las dimensiones. C) Los indicadores. D) La formación de los índices.

A) En la imagen inicial el o los autores tienen lo que puede calificarse como una idea más o menos vaga del concepto. A través de estudios de la bibliografía existente, tanto de la dedicada a la teoría como a la investigación, los autores han llegado a una imagen más clara del concepto, entendiendo por imagen más clara la posibilidad de poder realizar la segunda etapa, esto es:

B) La descomposición del concepto original en las dimensiones que lo componen. Antes de ver lo que trataron Adorno y colaboradores, damos un ejemplo de otro tipo: la inteligencia es un concepto complejo, cuyas dimensiones son inteligencia verbal, espacial, manual, numérica, abstracta, etc.; la participación política es otro concepto complejo, cuyas dimensiones pueden ser la forma de participación, los valores frente al sistema y actitudes, etcétera.

Adorno y colaboradores trabajaron cuatro dimensiones de la ideología general del autoritarismo. Éstas son: 1) el antisemitismo; 2) el etnocentrismo; 3) el conservadurismo político-económico, y 4) las tendencias antidemocráticas implícitas.

En cada una de estas dimensiones Adorno y colaboradores llegaron a establecer nuevas dimensiones (o subdimensiones, si queremos seguir nuestro razonamiento). No analizaremos todas las dimensiones, pero tomemos una como ejemplo: la más conocida

es la de las tendencias antidemocráticas implícitas. Este concepto de tendencias antidemocráticas implícitas incluye 9 dimensiones (o subdimensiones) que son las siguientes:

- 1) Convencionalismo (*conventionalism*).
- 2) Sumisión autoritaria (*authoritarian submission*).
- 3) Agresión autoritaria (*authoritarian aggression*).
- 4) Anti-intrasección (*anti-intrasection*).
- 5) Superstición y estereotipia (*superstition and stereotypy*).
- 6) Poder y dureza (*power and "toughness"*).
- 7) Destructividad y cinismo (*destructiveness and cynicism*).
- 8) Proyectividad (*projectivity*).
- 9) Sexo (*sex*).

C) El paso siguiente consiste en encontrar los indicadores para cada una de estas dimensiones. Y éstas son las preguntas concretas que se les hicieron a los sujetos para que respondieran según su grado de acuerdo o desacuerdo. Por ejemplo, la dimensión *convencionalismo* operó con algunos de los siguientes indicadores:

Ítem 1: "La obediencia y el respeto por la autoridad son las principales virtudes que debemos enseñar a nuestros niños".

Ítem 37: "Si la gente hablara menos y trabajara más, todos estaríamos mejor".

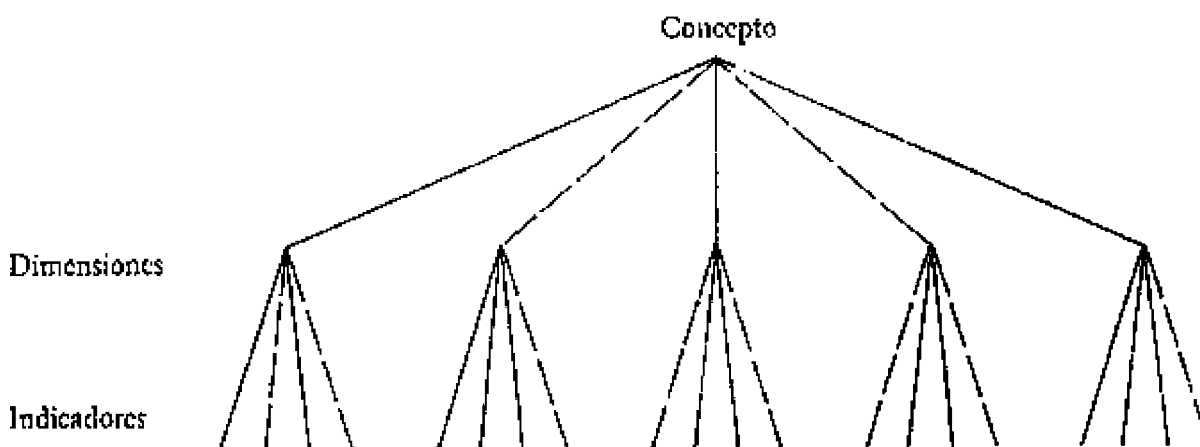
Ítem 9: "Cuando uno tiene problemas o preocupaciones, es mejor no pensar en ellos y ocuparse de cosas más agradables".

Ítem 41: "El comerciante y el industrial son mucho más importantes para la sociedad que el artista y el profesor".

Una pequeña acotación al respecto es que no nos parece correcto el procedimiento empleado por los autores de utilizar los mismos indicadores para una serie de dimensiones. Por ejemplo, el ítem: "Si la gente hablara menos y trabajara más, todos estaríamos mejor" aparece como indicador de convencionalismo, agresividad autoritaria y anti-intrasección. Si bien es aceptable que aquellas dimensiones estén intercorrelacionadas, esto no es demostrable utilizando este procedimiento, ya que el hecho de estar utilizando los mismos ítems hace que la correlación (si es que la hay) aparezca exagerada.¹⁷ Además, esto es un caso claro de tautología.

D) El paso final es recomponer el concepto original, uniendo las partes. Esta unión de partes es llamada aquí formación de índice. En el caso de los autores que tratamos, el concepto tendencias antidemocráticas implícitas es recompuesto a través de una escala (que es una forma particular de índice) bastante conocida como Escala F.

En un diagrama, los pasos a seguir en la operacionalización serían



Esta gráfica estaría indicando el paso del concepto a los indicadores. El paso de “retorno” de los indicadores al “nuevo” concepto (variable) puede tener mayor o menor grado de refinamiento, dependiendo de los objetivos de la investigación y/o del grado de precisión que se desee.

Este paso final de recomposición del concepto en un índice es la parte final de operacionalización del concepto en variable.

La palabra *variable* es entendida, en un sentido bastante amplio, como algo que “varía”, e incluye lo que podría denominarse atributo (nivel de medición nominal, donde la variación se da simplemente como un problema de distinción de clase), a cuasi variables (mediciones de nivel ordinal, donde además de distinciones de clase se hacen distinciones de grado en términos de rango), o de variables propiamente tales (niveles de medición intervalar y racional).

Índices

Los índices son síntesis de varias expresiones numéricas, y en nuestro caso particular el índice es la reconstitución de un concepto original que ha sido dimensionalizado y a cada dimensión se le han asignado diversos indicadores. Nos referiremos aquí principalmente a los índices sumatorios simples y a los índices ponderados. La adjudicación de índices simples comprende una serie de operaciones matemáticas relativamente arbitrarias, especialmente si se comparan con las escalas (ver capítulo sobre “Escalas para la medición de actitudes”).

Utilizamos índices sumatorios simples: *a)* cuando las mediciones son el resultado de una combinación de puntajes adjudicados arbitrariamente; por ejemplo, a las respuestas “de acuerdo” les adjudicamos un puntaje 2, y a las respuestas “en desacuerdo”, un puntaje 0; *b)* cuando los puntajes arbitrarios son indicadores de algunas dimensiones que queremos expresar en una cifra única; *c)* cuando el concepto no es escalable, esto es, cuando no podemos asumir la existencia de un continuo subyacente.

Construcción de índices. La construcción de un índice comienza a partir de algunos indicadores que pueden ser respuestas a alternativas (acuerdo-desacuerdo; sí-no; empleado-desempleado; urbano-rural; primaria incompleta-universitaria completa-

universitaria incompleta; valores porcentuales de urbanización, etcétera).

En la mayoría de los casos conviene que los valores asignados a cada una de las alternativas sean números enteros y positivos, sobre todo en aquellos casos donde el número asignado simplemente significa la asignación a una clase (medición nominal) o un orden de rango (medición ordinal). Conviene asegurarse también de que los valores del índice sigan la dirección de la variable, esto es, los números altos deben expresar los valores altos en la variable. Estos procedimientos nos ayudan en el proceso de interpretación, sobre todo cuando se calculan medidas de asociación y de correlación.

La amplitud del índice depende tanto de la cantidad de preguntas o de indicadores como de la cantidad de alternativas en cada uno de los indicadores. El valor mínimo del índice (si todos los indicadores tienen alternativas que expresen los valores mínimos con el número 0) será 0. Supongamos que tenemos un índice compuesto de 3 indicadores, cuyas alternativas de respuestas y pesos son, respectivamente:

Educación

Sin educación	0
Primaria incompleta	1
Primaria completa	2
Secundaria completa o incompleta	3
Universitaria completa o incompleta	4

Ocupación

Trabajadores manuales no especializados, asalariados y por cuenta propia	0
Trabajadores manuales especializados, asalariados y por cuenta propia	1
Personal de supervisión	2
Empleados de administración pública, privada y de comercio de grado II y comerciantes minoristas	3
Empleados de administración pública, privada y de comercio de grado I y propietarios medianos	4
Jefes menores e intermedios de administración pública, privada y comercio	5
Grandes empresarios y altos ejecutivos	6

Ingreso

Menos de 1 000 pesos	0
de 1 000 a 2 999 pesos	1
de 3 000 a 4 999 pesos	2
de 5 000 a 6 999 pesos	3
de 7 000 a 8 999 pesos	4
de 9 000 a 10 999 pesos	5
de 11 000 a 12 999 pesos	6
de 13 000 a 14 999 pesos	7
15 000 o más	8

Al índice resultante le llamamos de *status* socioeconómico, y tendrá un valor mínimo de 0 y un valor máximo de 18. Los valores extremos del índice (0 y 18) son el resultado de combinaciones únicas de indicadores (sin educación —no manual— con ingresos inferiores a 1 000; y universitaria completa o incompleta —grandes empresarios— con ingresos de 15 000 pesos o más). Los valores intermedios pueden ser el resultado de pautas diferentes. Por ejemplo, un puntaje 10 puede ser el resultado de varias combinaciones:

- Educación universitaria–personal de supervisión–ingreso entre 7 000 y 8 999.
- Educación primaria completa–pequeño propietario de comercio–ingresos entre 9 000 y 10 999.
- Educación secundaria completa–empleado de administración pública de grado II–ingresos entre 7 000 y 8 999.

El índice de *status* socioeconómico presentado como ejemplo tiene un problema adicional relativo al peso que tiene cada indicador. El peso máximo está en ingreso (representa 45% del valor total del índice); seguido por ocupación (33% del valor), y finalmente por educación (22%).

A menos que existan razones teóricas importantes que justifiquen la asignación diferencial de pesos, conviene asignar a los distintos indicadores la misma cantidad de alternativas (en el ejemplo presentado se pueden aumentar las alternativas de respuesta en educación).

Las alternativas de respuesta pueden ser construidas a partir de criterios teóricos o de criterios empíricos. Las alternativas resultantes de criterios teóricos van a ser tantas cuantas la teoría requiera, asignando los pesos según indicaciones que se desprenden de la conceptualización teórica. En el caso de asignación de alternativas según criterios empíricos, es posible discriminar, por ejemplo en los extremos de la distribución, según la distribución real de la población que se está estudiando. Tal puede ser el caso en algunas investigaciones que incluyen muestras representativas de la población y donde se hace necesario (en educación, por ejemplo) discriminar en los niveles bajos, por ejemplo:

Sin escolaridad	0
Hasta 2° grado de primaria	1
Hasta 4° grado de primaria	2
Más de 4° grado pero no completa ..	3

Primaria completa	4
Secundaria incompleta	5
Secundaria completa	6
Universitaria completa o incompleta	7

Es posible utilizar asimismo los valores parámetros para determinar los cortes. Por ejemplo, en ingreso, es posible tomar el promedio de la población y las desviaciones estándar para determinar diferentes cortes. Hay numerosas alternativas que dependerán del tipo de investigación que se está realizando, el tipo de variable, los conocimientos en el área, etcétera.

Los índices ponderados. Si por razones teóricas deseamos que los indicadores tengan pesos diferenciales, la asignación de pesos a las alternativas de respuesta en los indicadores que se desea que tengan valores mayores puede realizarse de la siguiente forma. Supóngase que el índice total esté compuesto de tres indicadores que tengan la misma cantidad de alternativas de respuesta, pero donde queremos adjudicar al indicador A un peso 2 (ponderación), mientras que los indicadores B y C van a tener un peso 1. La situación en cuanto al peso definitivo será entonces:

<i>Indicadores</i>	<i>Valores originales</i>			<i>Ponderación</i>			<i>Valores finales</i>		
A	0	1	2	2	2	2	0	2	4
B	0	1	2	1	1	1	0	1	2
C	0	1	2	1	1	1	0	1	2

Como puede observarse, la adjudicación de pesos se hace para cada valor original, lo que permite al investigador realizar una ponderación, si así lo desea, no sobre todos los valores originales, sino sobre alguna(s) alternativa(s) de respuesta en particular. Por ejemplo:

<i>Indicador</i>	<i>Valores originales</i>			<i>Ponderación</i>			<i>Valor final</i>		
D	0	1	2	1	1	3	0	1	6

Existen, por supuesto, otras alternativas de ponderación, distintas de la multiplicación simple, o de la suma, pudiéndose realizar algún otro tipo de operación matemática.

Cómo reducir un índice. Muchas veces los valores o el puntaje total de un índice pueden ser de tal amplitud que dificulten la tarea del análisis (por ejemplo, en casos de tabulaciones cruzadas); en tal caso el investigador pudiera estar interesado en reducir los valores de índice a tres o cuatro categorías, por ejemplo: alto; medio-alto; medio, y bajo. Dónde hay que realizar los cortes dependerá nuevamente de consideraciones teóricas, o

de criterios empíricos o de una combinación de ambos. Si es posible suponer que la variable que representa el índice es continua, es posible realizar los cortes de manera tal que se puedan garantizar suficientes casos en cada una de las celdas de la matriz de datos con las que se va a realizar el análisis. O es posible seguir los “cortes naturales” de la distribución de la variable, cuando se tienen las distribuciones empíricas. Supóngase la siguiente distribución de los puntajes de un índice de religiosidad:

<i>Puntajes</i>	<i>Frecuencias</i>	<i>Puntajes</i>	<i>Frecuencias</i>
0	2	12	14
1	3	13	6
2	3	14	12
3	7	15	12
4	6	16	12
5	10	17	8
6	8	18	8
7	10	19	7
8	9	20	4
9	6	21	4
10	11		
11	13	Total de casos	175

Si el investigador desea reducir, digamos, a 4 categorías la variable, existen varias posibilidades: es posible tomar los valores de índice, independientemente de las frecuencias en cada uno de ellos, y realizar los cortes, digamos, en las siguientes partes:

Bajo:	de 0 a 5
Medio:	de 6 a 10
Medio-Alto:	de 11 a 15
Alto:	de 16 a 21

En cuyo caso la distribución de frecuencia en cada una de las categorías sería ahora:

Bajo	31
Medio	44
Medio-Alto	57
Alto	43

Una alternativa distinta sería seguir los “cortes naturales” de la distribución, proponiendo los siguientes cortes:

<i>Valor</i>	<i>Puntajes</i>	<i>Frecuencia</i>
Bajo	0 a 4	21
Medio-Bajo	5 a 9	43
Medio	10 a 12	38
Medio-Alto	13 a 16	42
Alto	17 a 21	31

*La utilización de puntajes en un indicador, para otro indicador
en el que no se tiene respuesta*

Sucede frecuentemente que algunos de los indicadores no tienen respuesta, porque el sujeto no ha respondido a la pregunta o porque el entrevistador no ha registrado la respuesta o algo semejante. Supóngase que tenemos datos para la educación y la ocupación del sujeto, pero no hay respuesta para su ingreso. Queremos construir un índice de *status* en base a los tres indicadores y no queremos perder un caso. Es posible, asumiendo que existe intercorrelación entre los indicadores, utilizar las respuestas del sujeto en educación y ocupación para estimar su valor en ingreso. Los métodos más refinados incluirían regresiones para poder estimar el valor. Aquí señalaremos procedimientos algo más arbitrarios, y más inexactos.

Si el sujeto tiene un puntaje 2 en educación y 2 en ocupación, y si las variables están intercorrelacionadas, puede asumirse que el puntaje total será de 6 (es decir, asumimos que tendrá un valor 2 en ingreso). Para el caso de mayor número de indicadores y donde éstos se distribuyen en forma menos uniforme, es posible calcular valores promedios y adjudicarlos como el valor correspondiente a la variable en la que no se dispone del dato.

Cuando se calculan intercorrelaciones entre indicadores y no queremos perder casos, también es posible adjudicarle al sujeto los valores promedio de su grupo. Estas situaciones deben evitarse cuando la cantidad de casos sin datos supera un 10% de la muestra.

Ventajas y desventajas. Los índices pueden construirse en base a datos primarios o datos secundarios, presentando las siguientes ventajas: *a)* son fáciles de construir, *b)* las operaciones matemáticas son elementales, *c)* reflejan cambios importantes en las opiniones de sujetos, *d)* permiten evaluar mejor la situación, en la medida en que un solo indicador puede llevar a errores de diagnóstico.

Las principales desventajas: *a)* representan solamente una ordenación de los sujetos, *b)* no reflejan cambios pequeños de opinión, actitud o posición, ya que pautas distintas en conceptos multidimensionales pueden recibir puntajes similares.

A pesar de estas desventajas los índices son una herramienta útil para la investigación, especialmente cuando no se cuenta con los recursos para intentar mediciones más sofisticadas.

APÉNDICE

Consideramos conveniente presentar algunos ejemplos de taxonomías en tres disciplinas (botánica, psicología y sociología), por la utilidad que puedan tener con fines de orientación del lector con estas áreas.

La taxonomía en botánica es presentada a un lector que seguramente está familiarizado con los términos debido a sus estudios en la escuela media, y su propósito es facilitar la comprensión de la parte formal de las taxonomías en disciplinas como la psicología y la sociología, donde no existe aceptación universal de las propuestas. El lector notará que, si bien la taxonomía presentada como ejemplo para la psicopatología

es relativamente completa, existen problemas en la separación tanto en las delimitaciones de normalidad y anormalidad como en los diferentes síndromes que constituyen una enfermedad en particular. Esto es debido tanto al carácter dinámico de la enfermedad cuanto al estado teórico actual de la psicología. Lo importante a distinguir aquí, bien a nivel de la taxonomía psicológica, bien de la sociológica, es que no se atribuye en la taxonomía la explicación de por qué se da este tipo de síntoma o síndrome y no otro. Las teorías psicológicas sobre la enfermedad son abundantes, igual que las teorías en sociología.¹⁸

I) Lo que se conoce como *botánica sistemática* es un ejemplo excelente de taxonomía. Se trata de ordenar los vegetales agrupándolos en categorías que incluyen no solamente las plantas que existen en la actualidad, sino que también llega a incluir vegetales desarrollados en otras épocas. Las especies vegetales son designadas con nombres precisos e invariables que permiten la comunicación entre botánicos de todo el mundo. La taxonomía que presentamos es la de Linneo (1707-1778). Antes de Linneo, las clasificaciones de las plantas estaban fundadas en criterios prácticos o empíricos, poco sistematizados y bastante confusos, en los que se entremezclaban anécdotas, opiniones personales, antropocentrismos, etc. Linneo propone un sistema clasificatorio coherente, según la disposición de los órganos sexuales de las plantas, criterio que va a permitir clasificar cualquier planta en términos de familia, género y especie.

Para la botánica contemporánea, la clasificación de Linneo resulta algo simple, existiendo taxonomías más complejas basadas en características fisiológicas y morfológicas de las plantas, en las que éstas se ordenan, por ejemplo, según la división del trabajo fisiológico (sistemas de Braum, de Eichler, de Wettstein, sistema natural de Engler, etcétera).

CUADRO 1. *Taxonomía de Linneo; sistema sexual para la clasificación de las plantas*

Phanerogamia (plantas con flores)	Flores hermafroditas	Estambres libres	Estambres iguales	12 o menos estambres	1 estam. Clase 2 " " 3 " " 4 " " 5 " " 6 " " 7 " " 8 " " 9 " " 10 " " 12 " "	1 (Monandria) 2 (Diandria) 3 (Triandria) 4 (Tetrandria) 5 (Pentandria) 6 (Hexandria) 7 (Heptandria) 8 (Ocrandria) 9 (Enneandria) 10 (Decandria) 11 (Dodecandria)															
							Más de 12 estambres	Unidos al caliz " al receptác.	12 (Icosandria)	13 (Polyandria)											
									Dos o más largos	"	14 (Didynamia)	15 (Tetradynamia)									
							Estambres desiguales	En un haz	En dos haces	En varios haces	Por las antenas	Al pistilo	16 (Monadelphia) 17 (Diadelphia) 18 (Polyadelphia) 19 (Syngenesia) 20 (Gynandria)								
														Flores uni- sexuales	Monoicas	Dioicas	21 (Monoecia) 22 (Dioecia)				
																		Flores polígamas	23 (Polygamia)		
																				Cryptogamia	24 (Cryptogamia)

FUENTE: M. Ruiz O., D. Nieto R. e Ignacio Lario R.: *Tratado elemental de botánica*; Porrúa, México, 1950.

Las clasificaciones de la biología están montadas según principios lógicos similares. El sistema sexual o la taxonomía de Linneo aparece resumida en el cuadro 1.

II) *Taxonomía en las ciencias sociales.* Las taxonomías en las ciencias sociales, contrariamente a las existentes en las ciencias físicas y biológicas, no llegan a alcanzar ni el grado de completitud de aquéllas, ni tampoco niveles de aceptación universal.

Es quizá en psicología donde se ofrecen las taxonomías más completas, aunque existen distintos desacuerdos en cuanto a algunos supuestos básicos. Por ello las taxonomías en existencia se formulan según distintos modelos conceptuales. Las taxonomías en psicopatología, por ejemplo, se organizan por lo común alrededor de la dicotomía médica salud-enfermedad, de donde se establecen criterios y clasificaciones basados en la presencia de ciertos tipos de conducta, llamados síntomas, agrupados en síndromes que conjugan determinados tipos de desorden psicológico.

Los criterios para la agrupación de síndromes están basados en diagnósticos psiquiátricos, desajustes sociales, inventarios psicológicos objetivos, infelicidad subjetiva, presencia en hospitales mentales, etcétera.

Los criterios de normalidad y anormalidad a partir de los cuales se establecen las categorías y clasificaciones varían, por lo general, en tres tipos: ¹⁹ i) la normalidad como un concepto estadístico (la normalidad como las tendencias más frecuentes en la población); ii) la normalidad como un concepto ideal de salud mental, y iii) la anormalidad como la presencia de ciertos tipos de conducta que llevan al sujeto a la miseria personal o a la incapacidad para el manejo de sus propios asuntos. La

clasificación o taxonomía que vamos a presentar se basa principalmente en una serie de síntomas agrupados en síndromes, señalando para cada uno de los casos indicadores en las áreas somática, afectiva, cognitiva y motora.

I) *Normalidad*. La idea de *salud mental* o de *normalidad*, de Jahoda,²⁰ entremezcla conceptos de salud mental como estado ideal, así como la presencia de ciertos atributos positivos; los criterios son: i) *auto-comprensión (self-insight)*, esto es, predominancia de lo consciente sobre lo inconsciente; ii) *balance de fuerzas psíquicas*; iii) *auto-actualización (self-actualization)*; iv) *resistencia a las tensiones (stress)*; v) *autonomía*; vi) *competencia*, y vii) *percepción de la realidad*.

II) *El dominio de la psicopatología* es agrupado tradicionalmente en cuatro grandes áreas:²¹ *neurosis, psicosis, trastornos psicosomáticos y trastornos de la conducta*.

CUADRO 2. *Taxonomía en psicología-psicopatología*

A) <i>Normalidad</i> : salud mental y presencia de atributos positivos.					
B) <i>Anormalidad (psicopatología)</i>	1) Neurosis	{	a) Reacciones de ansiedad y fobias		
			b) Histeria		
			c) Reacciones obsesivo-compulsivas		
			d) Reacciones de depresión		
			e) Fatiga — hipocondriasis		
	2) Psicosis	{	a) Afectivas	{	i) manía
					ii) depresión
			b) Esquizofrenia	{	i) Simple
					ii) Hebefrenia
					iii) Catatónica
					iv) Paranoide
	3) Trastornos en la conducta	{	a) Psicopatía		
			i) Drogas		
			ii) Alcoholismo, etc.		
			c) Desviaciones sexuales	{	i) homosexualidad
					ii) sadismo-masochismo
					iii) voyeurismo
					iv) fetichismo
					v) exhibicionismo
					vi) travestismo, etc.
	4) Trastornos psicosomáticos	{	a) Úlcera péptica		
			b) Neurodermatosis		
			c) Hipertensión esencial		
			d) Asma bronquial, etc.		

A) *Las neurosis* comprenden una serie de síntomas, que no constituyen una clara ruptura con la realidad, que representan cierta permanencia en el individuo o que concurren periódicamente sobre lapsos extendidos de tiempo. Buss incluye la siguiente lista de síntomas neuróticos que hacen que el individuo tienda a ser sobreinhibido, ansioso o sobrepreocupado con culpa (*guilt-ridden*); esas tendencias obstruyen la acción directa en la resolución de los problemas en la vida cotidiana; consecuentemente, la persona neurótica tiende a enfrentar la necesidad de elección con vacilación, la necesidad de dirección con extravío, y la necesidad de acción enérgica con timidez. i) Excesivo olvido, pensamientos obsesivos y rituales compulsivos —intentos autoderrotantes (*self-defeating*) para hacer frente a situaciones que producen miedo—. ii) Depresión y fatiga

—residuos psicológicos de tensiones prolongadas—.

Tradicionalmente los síntomas neuróticos han sido agrupados en 5 clases de *síndromes*:

a) *Reacción de ansiedad y fobias*. Caracterizada por excesivo miedo, cuyos estímulos son desconocidos en el caso de los primeros y específicos en las fobias. Los síntomas de *free-floating* ansiedad se agrupan en cuatro sistemas de reacciones:

1) *Somáticos*: centelleos, sudor, resequedad en la boca, opresión en el pecho, taquicardia, incremento en la presión sanguínea, dolores de cabeza, sensación de debilidad, trastornos intestinales, etcétera.

2) *Afectivos*: agitación, pánico, depresión, irritabilidad.

3) *Cognitivos*: preocupación, pesadillas, temor, distracción, olvido.

4) *Motores*: tensión muscular, temblores, falta de coordinación, “congelamiento”, sobresalto.

Las *fobias* más comunes son hacia: animales, multitudes, gérmenes, oscuridad, venenos, lugares elevados, tormentas, lugares cerrados, etcétera.

b) *Histeria*. Incluye una gran variedad de síntomas que superficialmente guardan parecido con algunas enfermedades orgánicas, pero en los cuales no existen bases orgánicas. Los síntomas de la histeria (aunque el histérico individual presenta usualmente síntomas en una o dos esferas), según los 4 sistemas de reacciones mencionados más arriba, son:

1) *Somáticos*: dolor, debilidad, enfermedades pseudoorgánicas, desmayos.

2) *Afectivos*: indiferencia.

3) *Cognitivos*: pérdida de visión, pérdida de sentido del oído, pérdida del sentido del tacto, pérdida de memoria, sonambulismo.

4) *Motores*: pérdida del habla, de la locomoción, tics, calambres musculares, ataques.

c) *Reacciones obsesivas-compulsivas*. La obsesión y la compulsión tienden a concurrir en el mismo paciente, existiendo desacuerdo en cuál de las dos tiende a ser más predominante; de ahí que se usen a veces los términos neurosis obsesiva o neurosis compulsiva. La obsesión es una preocupación excesiva por ciertos temas, con la exclusión de cualquier otro. La sintomatología tiende a darse en dos áreas: 1) *dudas obsesivas*, que colocan al individuo en un estado perenne de indecisión; o 2) *pensamientos obsesivos*, que involucran acciones o tentaciones prohibitivas, de tipo agresivo o sexual; estos impulsos son sentidos como exteriores, con sensación concomitante de horror y repulsión por parte del individuo obsesivo.

Las *compulsiones* son secuencias de conducta que impelen al individuo a llevarlas a cabo; de no realizarse el ritual, se genera ansiedad. Las compulsiones más comunes son las de repetición y simetría y exagerada atención al detalle. En las primeras el individuo repite la secuencia de conducta una y otra vez. En las compulsiones simétricas, la última parte de la secuencia de conducta balancea la primera parte.

d) *Depresión neurótica*. Cae dentro del tipo de síndromes productos de residuos psicológicos de tensiones. En este sentido la depresión neurótica puede ser vista como una reacción a la pérdida o al fracaso: fracaso en la resolución de problemas importantes

de la vida cotidiana; fracaso para enfrentar el conflicto, o pérdida de una relación importante. Los síntomas se dan fundamentalmente en las áreas afectivas y motoras:

1) *Afectivos*: melancolía, pesimismo, apatía, autodepresión.

2) *Motores*: lentitud en el habla; lentitud de movimiento.

e) *Fatiga hipocondriasis*. Combina dos clases de síntomas que a veces se mantienen separados: cansancio y sobrepreocupación con la salud. Los síntomas tienden a concurrir en dos esferas: la afectiva y la somática:

1) *Somáticos*: dolores y padecimientos; falta de energía.

2) *Afectivos*: aburrimiento, irritabilidad.

B) *Las psicosis* comportan síntomas más serios y desviados que los neuróticos: i) Alucinaciones: perturbaciones en percibir la realidad. ii) Delusiones: perturbaciones en interpretar la realidad. iii) Pérdida de la capacidad de asociación, habla fragmentada e incoherente: perturbaciones en el pensamiento. iv) Melancolía prolongada o júbilo: perturbaciones en el humor. v) Aislamiento y retiro de otras personas: perturbaciones en la interacción personal.

Siguiendo a Buss (*op. cit.*), clasificamos los síntomas de psicosis en dos grandes grupos: psicosis afectivas y esquizofrenia.

a) *Las psicosis afectivas* incluyen los síntomas en los cuales el humor o estado de ánimo es predominante. La división en síndromes se hace comúnmente entre manía y depresión, que con frecuencia se dan como alternantes en el mismo paciente.

1) *Manías*. Los síntomas se dan principalmente en las esferas motoras y afectivas. Los niveles de actividad varían desde el exceso de energía en júbilo al retardo motor en la melancolía: i) *síntomas afectivos*: humor jovial que se expande a toda la esfera de actividades, superoptimismo, euforia, grandiosidad, alta autoestima y baja autocomprensión (*self-insight*), pomposidad, alto sentido gregario, excesos en la sociabilidad que lleva a la intrusión y a la molestia. ii) *Síntomas motores*: movimientos físicos excesivos, euforia motora, falta de inhibición y control, impulsividad e impetuosidad, excitamiento y excitabilidad, a veces conducta violenta (agresiva o sexual). iii) *Cognitivos*: atención intensa pero breve, asociaciones tangenciales leves, presencia de algunas ideas ilógicas, rapidez en el habla, fantasías egocéntricas, delusiones de grandeza o de persecución leves y libres de hostilidad.

2) *Depresión*. Es más compleja que la manía. A veces se clasifica en subtipos: simples, agitadas, involucional y *post-partum*. Los síntomas más importantes se dan en las esferas afectiva y motora, representando la contraparte de los que se manifiestan en la manía: i) *afectivos*: sentimientos de culpa, remordimientos, pesar, aburrimiento y desinterés, falta de esperanza, pesimismo, desalentamiento. ii) *Motores*: letargia, indolencia, falta de iniciativa, autoagresión que a veces lleva al suicidio. En los depresivos agitados, en los que se dan altos contenidos de ansiedad, hay algún tipo de actividad. iii) *Somáticos*: bajo tono digestivo y en general de los procesos gastrointestinales, pérdida de apetito y de peso, insomnio; constipaciones, sentimiento de presiones en la cabeza y pecho. iv) *Cognitivos*: depersonalización, inhabilidad para

sentir, amar, para experimentar placer. Delusiones de carácter leve y circunscritas a tres áreas: culpa (reclamo de responsabilidad sobre todo accidente, pérdidas, daños, etc.); nihilismo (el paciente o el mundo están en un curso de total destrucción del que no hay escape) y procesos corporales (desintegración del cuerpo).

b) *Esquizofrenia*. Tradicionalmente se delinean 4 tipos de esquizofrenia:

1) *Esquizofrenia simple*: solapado alejamiento en la interacción social, pobreza de pensamiento (sin alucinaciones o delusiones), profunda apatía y falta de respuesta emocional.

2) *Hebefrenia*: conducta infantil y necia; deterioro en los procesos del pensamiento; habla incoherente, neologismos, asociaciones extrañas; alucinaciones, y a veces delusiones.

3) *Catatónica*: extremos en el nivel motor de actividad (rigidez o excitación), alucinaciones.

4) *Paranoide*: incluye solamente síntomas cognitivos; presencia de delusiones y alucinaciones.

Los síntomas más importantes de la esquizofrenia se producen en la esfera cognitiva, particularmente como un resultado de alteraciones en la percepción de estímulos; en general conducen a la depersonalización, y son: perturbaciones en el lenguaje y en el pensamiento, perturbaciones en el sentido de sí mismo (particularmente en la imagen corporal), autismo, alucinaciones y delusiones.

En el curso de la esquizofrenia es posible delinear 4 estadios sucesivos. Primer estadio, *ansiedad* (comienza con pánico ante los síntomas, seguido por intuición de tipo psicótico y proliferación de nuevos síntomas); el segundo estadio es llamado *avanzado* (el paciente acepta los desórdenes y deviene apatético, deterioro en el área afectiva y retiro de las áreas de interacción social); en el tercer estadio, o *preterminal*, los síntomas se agravan al punto que es difícil distinguir distintos tipos de esquizofrenia; en el cuarto estadio, o *terminal*, el paciente aparece dominado por hábitos primitivos. Generalmente los pacientes pueden no seguir el curso completo de la enfermedad, sino permanecer ya en el segundo estadio, o alternando periodos esquizofrénicos con periodos relativamente libres de síntomas.

C) *Desórdenes psicosomáticos*. Presentan un grupo de enfermedades que involucran trastornos orgánicos que están indistintamente asociados a perturbaciones emocionales y que involucran por lo general estructuras innervadas por el sistema nervioso autónomo: estómago (úlceras pépticas), piel (neurodermatosis), arterial-periféricas (hipertensión esencial), bronquios (asma bronquial).

D) *Trastornos de la conducta*. Se refiere de ordinario a trastornos no tolerados por la sociedad. Por lo general se consideran tres clases:

a) *Psicopatías*. Se dan dos clases de síntomas: la conducta asocial o antisocial (irresponsabilidad, amoralidad), y la inmadurez o falta de socialización (ausencia de culpa o vergüenza). Sintomatología:

1) Conducta orientada hacia las conmociones y desprecio por las convenciones (trata a las personas como objetos).

2) Inhabilidad para controlar impulsos o posponer gratificaciones.

3) Rechazo de la autoridad y la disciplina.

4) Juicio pobre sobre conductas, pero buenos juicios acerca de situaciones abstractas.

5) Falla en alterar o modificar una conducta que ha sido castigada (falta de sentimiento de culpa).

6) Mentiroso patológico.

7) Conducta asocial y antisocial.

Características de personalidad: *i)* En las relaciones personales, incapacidad para el amor o la amistad intensa, centrado en sí mismo, egoísta, falta de empatía, usa a las personas como objeto. *ii)* *Adicciones*: alcoholismo, drogas, etc. *iii)* *Desviaciones sexuales*: incluye homosexualidad; relación sexual con animales; sadismo-masochismo; voyeurismo y fetichismo; exhibicionismo; travestismo.

En sociología las taxonomías son menos precisas que en psicología y no existe un sistema universal para el diagnóstico sociológico aceptado.²² Presentamos aquí la tipología de Gino Germani,²³ en la que el autor, aclarando que se trata de una extrema simplificación, señala que los dos tipos de sociedades representan los extremos de un continuo pluridimensional. Los dos tipos de sociedades se describen en función de las modificaciones que se producen en tres principios básicos de la estructura social: 1) El tipo de acción social. 2) La actitud hacia el cambio. 3) El grado de especialización de las instituciones.

El proceso de transición se caracteriza como un proceso de secularización.

CUADRO 3. Dos tipos ideales contrapuestos: sociedad tradicional y sociedad industrial

Sectores	Sociedad tradicional	Sociedad industrial	
		Modelo "Liberal"	Transformaciones recientes
Principios básicos de la estructura social	Acción prescriptiva, institucionalización de la tradición, instituciones indiferenciadas	Acción electiva, institucionalización del cambio, especialización creciente de las instituciones	
Tipo de relaciones sociales características	Adscripción, particularismo, carácter difuso, carácter afectivo	Desempeño, universalismo, especificidad, neutralidad afectiva	
Tecnología	Utensilios manuales	Máquinas	
Energía:	Energía humana y animal	Energía proporcionada por motores "primarios"	Energía proporcionada por motores "secundarios". Energía atómica
Procedimientos de producción	"Artesanal" (unidad por unidad)	"Producción en serie". "Cinta de montaje"	"Automatización"
Actitudes:	Procedimientos tradicionales Se desalienta el cambio	Procedimientos "racionales", búsqueda del cambio	
Economía	"Economía de subsistencia" Producción para satisfacer necesidades concretas, de individuos o grupos concretos, en un nivel tradicional	Economía de producción para el cambio. Producción para satisfacer una "demanda", un público "comprador" abstracto Economía de mercado	
Rasgos generales		Hincapié sobre la producción	Economía de mercado y nuevas formas de regulación Hincapié sobre el consumo. Economía monetaria
	Economía "natural"	Carácter dinámico de la economía	
	Carácter estático de la economía	Interdependencia creciente, crisis económicas	
	Unidades productoras autosuficientes	La esfera de lo económico se especializa. Funcionalización y especialización de las actividades	
	La esfera de lo económico indiferenciada del sistema social	División del trabajo funcional según criterios racionales. Búsqueda de la eficacia. Importancia de la profesión	
	División del trabajo tradicional. Según status adscriptos, por sexo y edad	El trabajador no posee los instrumentos de producción	
	El trabajador posee los instrumentos de producción	Creciente importancia del capital fijo	
	Poca importancia del capital fijo (maquinaria, instalaciones, edificios, etc.)	Principio hedónico: mínimo esfuerzo	
Principios y hechos que rigen la organización económica	—Reciprocidad —Redistribución —Autoabastecimiento —Economía doméstica	Racionalización creciente: adecuación de medios a fines Fin "económico", "rentabilidad", diferenciación de otros fines "no económicos" Comercialización de los factores de la producción Especialización e interdependencia creciente (entre individuos, empresas, países)	
Formas de repartición del excedente económico	No estratificadas: repartición más igualitaria	Ganancias: "lucro racional" Libre competencia Empresas individuales Trabajo "libre" Precios regulados por el mercado No intervención del Estado	
	Estratificadas: repartición desigual Sociedades "feudales"	más desigual Limitaciones al principio de la ganancia; incluso del fin económico menos desigual	

CUADRO 3. (Continuación)

Sectores	Sociedad tradicional	Sociedad industrial	
		Modelo "Liberal"	Transformaciones recientes
Motivaciones y actitudes hacia la economía y el trabajo.	No estratificadas: faltan motivaciones especiales para la actividad económica Producir bienes de consumo, concretos, hasta cubrir el nivel de <i>subsistencia</i> fijado por la tradición. Trabajo directamente relacionado con la necesidad <i>Principio de la subsistencia</i> No hay espíritu de competición en lo económico Estratificadas: En general: mismos principios Clases superiores: lo económico considerado inferior: sólo consumen (no producen, debiendo los inferiores proporcionarles lo debido tradicionalmente) Clases inferiores: artesanos: "instinto artesanal"; siervos; aceptación pasiva, compulsión	Grandes empresas. Empresas directoriales, anónimas, mixtas, nacionalizadas Trabajo organizado. Sindicatos. Contratación colectiva Precios políticos Intervención, regulación, planificación, propiedad estatal Clases superiores (burguesía): <i>Homo economicus</i> Principio hedónico. Lucro: por medio de actividad productiva racional organizada: rentabilidad "Ascesis capitalista" Éxito económico identificado con éxito en la vida (cf. signo de salvación en el calvinismo) "Espíritu de empresa" Expansión, innovación	
Propiedad típica	Formas concretas de propiedad comunal y personal "primitivas" varia "feudales" tierra	Clases superiores: Aparecen nuevas motivaciones; disminuye deseo de "lucro racional". Competición por el <i>status</i>; poder y prestigio dentro de la empresa directorial; disminuye <i>ascesis</i> Clases inferiores: Competición por el <i>status</i>; deseo de seguridad Formas más abstractas de propiedad	La fábrica. Propiedad personal de capital mobiliario Propiedad de títulos representativos de capital o créditos. Participación en propiedad colectiva Empresa individual. Soc. de personas Empresa anónima. Sociedad de capital. Empresa directorial. Separación de la propiedad y del control privada mixta pública
Unidades económicas típicas	Son las mismas que corresponden a la organización social: Familia extensa Taller artesanal y actividad agrícola	Actividades secundarias	Actividades terciarias
Actividades económicas típicas	Actividades primarias		
Organización social	<i>Predominio de lo primario</i> Grupo de parentesco Grupo de edad Grupo de sexo Grupo local	<i>Predominio de lo secundario</i> clase social ocupación nacionalidad	
Tipo de <i>status</i>	En las sociedades estratificadas adquieren importancia: Castas, estamentos <i>Status</i> definidos por la edad, el	Multiplicidad de grupos "secundarios" <i>Status</i> definido por la clase, la ocupación, la pertenencia a grupos secundarios	

CUADRO 3. (Continuación)

Sectores	Sociedad tradicional	Sociedad industrial	
		Modelo "Liberal"	Transformaciones recientes
	sexo, el parentesco, la casta, el estamento		
	<i>Status adscriptos importantes</i>		<i>Status adquiridos Menos importantes</i>
Grupos primarios familia	Familia extensa u otras formas similares Funciones: biológicas, económicas (producción y consumo), educacionales; recreativas, religiosas, etcétera <i>Posición central del grupo de parentesco en la sociedad</i>	Familia nuclear Conyugal, aislada, inestable Funciones: biológicas; socialización del niño; ajuste emocional del adulto; económicas; consumo únicamente	
Grupo local	Aldea-vecindario	Aldea, vecindario El grupo de parentesco menos importante. El suburbio	
Grupos secundarios	<i>Poco importantes o inexistentes</i>	<i>Muy importantes</i> Ocupaciones. Grupos educacionales Recreativos. Ideológicos	De intereses Asociaciones voluntarias
Religión	<i>Importante</i> Penetra en toda la vida social	<i>Menos importante</i> Toca una esfera especializada: la religión como aspecto separado de la vida <i>Laicización</i>	
Estratificación social	No estratificadas No se distinguen claramente capas sociales superpuestas Estratificadas "feudales"	Burguesía: propietarios de la industria y del comercio; profesiones liberales Obreros industriales y otros rurales: propietarios arrendatarios, dependientes sin tierra	
	Aristocracia (terratenientes nobles; militares, sacerdotes)	Disminuyen los propietarios de industria y de comercio Aumentan dirigentes de empresas, técnicos; <i>Empleados</i> Obreros industriales (varía la composición). Disminuyen las clases rurales: propietarios y dependientes	
	Clases inferiores: Hombres libres (artesanos, mercaderes, etc.); siervos, ligados a la tierra, etcétera <i>Cerradas</i> <i>Estáticas</i>	<i>Clases abiertas. En teoría: absoluta movilidad social. En práctica: diferentes grados de movilidad social</i>	
	Poca o nula movilidad social <i>Adscripción</i>	<i>Adquisición a través de la lucha competitiva</i> <i>Igualdad de oportunidad</i>	
Ideologías relativas a la estratificación social	Sociedades estratificadas: A cada uno según su <i>status</i> , según el lugar que le tocó en la vida <i>Se desalienta la movilidad. Se estimula la permanencia en la misma posición</i>	Compulsión a ascender socialmente No ascenso = fracaso Posición = fruto del esfuerzo personal Ascenso a través del enriquecimiento por medio del éxito en los negocios propios; o llegando a transformarse de "dependiente" (obrero, por ejemplo) en "por su propia cuenta", etcétera Aparece la motivación a ascender dentro del sistema de posiciones de la gran empresa directorial pública o privada. Motivos del prestigio y el poder	
Aspectos demográficos	Poca población "Alto potencial demográfico"	Extraordinario aumento de población "Transición demográfica"	
	Alta natalidad Alta mortalidad	"Bajo potencial demográfico"	
		Alta natalidad (bajando) Mortalidad (bajando) Tasas diferenciales Clases med.: baja natalidad Clases pop.: alta natalidad	Disminuyen las diferencias demográficas entre clases

CUADRO 3. (Continuación)

Sectores	Sociedad tradicional	Sociedad industrial	
		Modelo "Liberal"	Transformaciones recientes
	Población esencialmente rural	Población crecientemente urbana <i>La gran ciudad</i> <i>Megalópolis</i>	
Centro típico	La aldea El vecindario	Centralización	Ligera tendencia a la descentralización Destrucción de la aldea Importancia del suburbio como centro "local"
Movilidad ecológica	Baja o ninguna	Alta, con tendencia a crecer	
<i>Tipos de autoridad y control</i>	Tradicional (y formas carismáticas) Costumbre	Estado moderno Racional — burocrática <i>Leyes-reglamentos</i>	Sociedad masificada: aparecen formas carismáticas. Conformismo
<i>Caracteres generales de la sociedad, la cultura y la personalidad</i>	"Sociedad sagrada aislada" Todas las funciones tienden a permanecer indiferenciadas, dentro del sistema social (familia, economía, religión)	"Sociedad secular accesible"	
Grado de homogeneidad:	<i>Alta homogeneidad</i> Repugnancia a lo distinto. Intolerancia. Etnocentrismo	Máxima especialización — diferenciación de las funciones en esferas separadas	<i>Alta heterogeneidad</i> Aceptación de lo distinto:
	Extranjero = extraño = enemigo	Tolerancia espíritu liberal cosmopolitismo	
Grado de cambio:	"Lo antiguo = sagrado" Repugnancia al cambio. Dominio de la tradición y sus portadores: los ancianos	Exaltación de lo nuevo Búsqueda del cambio Dominio de la voluntad de transformación basada en principios de racionalización	
Grado de comunicación:	Poca o nula. Pocos contactos	Multiplicación de contactos,	
Accesibilidad social y ecológica	Aislamiento social y ecológico	creciente comunicación, creciente accesibilidad social y ecológica	
Formas de sociabilidad	"Primaria" (el vínculo familiar)	"secundaria" (lo funcional y lo anónimo)	
Relación con el grupo:	Individuo sumergido en el grupo Sentimiento de pertenencia	Forma normal: "individuación-liberación"	Forma patológica: Sentimiento de aislamiento. "Atomización"
Grado de libertad (psico-social)	Baja	Alta	Surgen problemas vinculados a la personalidad. "Miedo a la libertad"
Tipo de integración	Basada en la tradición, la conformidad, la estaticidad, la inmersión en el grupo. La alta homogeneidad, factor de integración	Basada en la interdependencia funcional; el reajuste autónomo y funcional de individuos "liberados", coexistencia de lo heterogéneo; adaptación adecuada al cambio; elección de valores por el individuo por medios racionales y por el ejercicio de su voluntad. Formas patológicas de integración: anomia, desintegración social	
Sistema de valores. Contenido	Varia: en general de carácter religioso, trascienden al individuo y su vida terrenal, o al grupo como verdadera o superior realidad La tradición, la sangre, la tierra, la divinidad	Afirmación del individuo como ente autónomo dotado de facultades racionales, capaz de dirigirse a sí mismo apoyándose en sus propias fuerzas. Hincapié en valores inmanentes al individuo y a su vida terrenal Afirmación de la razón, la voluntad, el cambio (el "progreso"), la libertad, la tolerancia Aparecen formas contrarias de valoración: "raza", "sangre", "nación", resurgimientos religiosos o nuevas formas de religiosidad; el "Estado", la "clase"; irracionalismo, etcétera	
Forma de aplicación de los valores	El sistema de valoración es único en cada sociedad (homogeneidad); se fija clara y detalladamente la conducta del individuo en las diferentes situaciones vitales; el individuo no debe elegir; no debe	Hay multiplicidad de valores y de criterios de valoración a menudo en conflicto entre sí. Los individuos deben elegir por medio de su voluntad y razón. Las situaciones que enfrentan son cambiantes y pueden no responder a las expectativas: los individuos deben realizar continuos ajustes; hay ambigüedad y contradicción	

CUADRO 3. (Conclusión)

Sectores	Sociedad tradicional	Sociedad industrial	
		Modelo "Liberal"	Transformaciones recientes
	interpretar: sus actitudes internalizadas responden de manera automática a las situaciones que se le presentan	Lo internalizado no deberían ser formas rígidas de comportarse, sino capacidad de adecuarse creativamente al cambio	
Tipo de personalidad	"Tradicional"	Dirigida desde adentro: "giroscopio" Dirigida desde afuera: "radar"	

III. MUESTREO

JORGE PADUA

“UNIVERSO” o “población” son palabras utilizadas técnicamente para referirse al conjunto total de elementos que constituyen un área de interés analítico. Lo que constituye la población total está delimitado, pues, por problemáticas de tipo teórico: si la referencia es a individuos humanos, el universo o población estará constituido por la población total de la humanidad; o por la población de un país, o la de un área determinada, etc., etc., según sea la definición del problema de investigación.

Los elementos que constituyen a una población, por supuesto, no tienen que ser necesariamente individuos humanos; uno puede referirse a naciones, grupos, edificios, animales, objetos físicos o elementos abstractos (tales como, por ejemplo, la “población de verso y reverso de una moneda” o de cualquier otra distribución binomial o multinomial).

Los parámetros poblacionales (o simplemente parámetros, o valores “verdaderos”) caracterizan las distribuciones en la población o universo; éstos pueden ser valores de ciertas distribuciones de variables aleatorias tales como la media aritmética o la desviación estándar.

Se denomina *muestra* un subconjunto del conjunto total que es el universo o población.

La teoría del muestreo tiene como propósito establecer los pasos o procedimientos a través de los cuales sea posible hacer generalizaciones sobre la población a partir de un subconjunto de la misma, con un grado mínimo de error. Sin embargo, como veremos más adelante, no toda muestra tiene como propósito “sacar conclusiones” acerca de la población; existen varios tipos de muestreo que se ocupan de selección de muestras para otros propósitos.

Los valores muestrales son los estadísticos computados a partir de las muestras, y con los cuales se buscará estimar los parámetros poblacionales.

Decíamos que no toda muestra o no toda investigación tiene como propósito obtener conclusiones acerca de la población. Según Galtung¹ una muestra debe satisfacer en general dos condiciones: 1) en ella debe ser posible poner a prueba hipótesis sustantivas, esto es, proposiciones acerca de relaciones entre variables, y 2) debe ser posible poner a prueba hipótesis de generalización —de la muestra al universo— sobre las proposiciones establecidas en la muestra.

La primera de las condiciones hace referencia al hecho de que la muestra sea lo suficientemente “buena” para permitirle al investigador extraer conclusiones en cuanto a las relaciones entre sus variables.

La segunda de las condiciones tiene que ver con la posibilidad de establecer generalizaciones; es decir, inferencias válidas, con un grado de incertidumbre conocido. Para esto es necesario que las muestras sean probabilísticas, ya que la determinación del grado de incertidumbre o de “confianza” que pueda atribuirse a las inferencias depende en sus cálculos de la teoría de las probabilidades.

Si bien las muestras aleatorias o probabilísticas (cuando son de tamaño grande) tienden a cumplir las dos condiciones, existen casos en los cuales el investigador puede estar más interesado en la teoría sustantiva que en la generalización. Dedicamos la última sección a los procedimientos para selección de muestra en este tipo de situaciones.

Siguiendo los propósitos de este manual, vamos a concentrarnos más en aspectos operacionales que en aspectos “teóricos” o “sustantivos”; es decir que no vamos a profundizar en la teoría estadística del muestreo ni tampoco en los problemas específicos que surgen en las investigaciones concretas y que tienen que ver con la relación entre la muestra y el tipo de hipótesis que el investigador quiere probar. Hay que señalar solamente que en algunas ocasiones la muestra en sí puede “desvirtuar” la investigación en el sentido en que los procedimientos muestrales pueden conducir a una desestructuración del universo; por ejemplo, cuando las unidades están conectadas por relaciones específicas de dominación, de interacción, de comunicación, etc. Nosotros vamos a tomar estos casos como un tipo especial de muestras predispuestas; es decir, muestras en donde la comprobación o refutación de la hipótesis pasa a ser el resultado de procedimientos de muestreo.

Puesto que vamos a concentrarnos en señalar los modos operacionales para obtener distintos tipos de muestras y no en la teoría estadística de las distribuciones muestrales, recomendamos al lector interesado algunos textos importantes en esa línea.² La literatura en el área es abundante, y la lista señalada es solamente una mínima selección.

DISTINTOS TIPOS DE MUESTRAS

El problema del muestreo surge cuando la población a estudiar es demasiado numerosa como para implicar costos en energía y dinero insuperables. Se trata entonces de seleccionar a un subconjunto que minimice esos costos al mismo tiempo que no se produzca pérdida de precisión.

La teoría del muestreo establece las condiciones mediante las cuales las unidades o las muestras son seleccionadas de manera tal que el subconjunto resultante (la muestra) contenga el mínimo de sesgos posibles.

Muestras predispuestas. Son aquellas que han sido seleccionadas de manera tal que la comprobación o la refutación de las hipótesis pasa a ser el resultado de procedimientos de muestreo. Por ejemplo, una muestra predispuesta sería aquella en la cual, para comprobar alguna hipótesis acerca del comportamiento político de la población total de votantes, utilizamos solamente individuos o elementos que, a partir de su pertenencia a un grupo específico, presentan peculiaridades que no son típicas de la población en su

conjunto.

CUADRO 1

		Procedimientos básicos	Ventajas		Desventajas		
			Técnicas	Económicas	Técnicas	Económicas	
Muestras probabilísticas (Todos los individuos o elementos tienen una probabilidad conocida de ser incluidos en la muestra)	Todas las combinaciones tienen igual probabilidad de darse en la muestra	Muestra simple al azar	1. Hacer una lista completa del universo 2. Asignar un número a cada individuo del universo 3. A través de una tabla de números aleatorios o procedimiento similar seleccionar un número de individuos que van a constituir la muestra	1. Igual probabilidad de inclusión en la muestra de todos los individuos 2. Permitir la generalización 3. Proporciona base para calcular el grado de dispersión entre las medidas de la muestra y del universo	Las que se derivan de no tomar a todo el universo	1. No provee un número suficiente de casos de grupos especiales 2. Puede haber distorsiones en cuanto a la representatividad	1. Alto costo monetario 2. Alto costo de tiempo
	Necesariamente no todas las combinaciones tienen igual probabilidad de darse en la muestra	Sistemática	1. Hacer una lista completa del universo 2. Seleccionar el primer individuo a través de un método aleatorio 3. Seleccionar cada k -ésimo individuo a partir del 1º seleccionado (por ejemplo, cada décimo individuo)	Idem que para la muestra simple al azar	1. Mayor facilidad en obtener la muestra	1. Si hay algún tipo de bias en la lista el muestreo puede resultar influido por él	Idem
		Estratificada	1. Dividir el universo en estratos internamente homogéneos 2. Seleccionar dentro de cada estrato los individuos de modo aleatorio 3. Las fracciones de muestra, en cada estrato, son proporcionales	1. Idem que para los dos anteriores 2. Garantiza la representatividad 3. Elimina los errores entre estratos	1. Idem que para la muestra sistemática	1. Puede no proveer un número suficiente de casos para estratos pequeños 2. Dificultad para determinar estratos homogéneos	1. Costo más alto que en el muestreo simple y de conglomerado 2. Alto costo de tiempo
			1. Idem, proporcional 2. Idem, proporcional 3. Las fracciones en cada estrato pueden ser distintas según las necesidades	1. Idem que para los dos anteriores 2. Posibilita mejor conocimiento de grupos pequeños en el universo	1. Si los estratos son homogéneos <i>ceteris paribus</i> , hace posible una muestra menor y mayor economía	1. Exige tratamientos estadísticos algo complejos 2. Dificultad para determinar estratos homogéneos	1. Costos más altos que en el muestreo de tipo cluster
		Conglomerados	1. Dividir el universo en diversos grupos o clusters 2. Seleccionar primero qué clusters deben constituir la muestra 3. Dentro de cada cluster seleccionar los individuos de la muestra de modo aleatorio	1. Idem que para los dos primeros	1. Ahorra dinero sobre todo porque permite la concentración de los entrevistados en áreas próximas. 2. Ahorra tiempo	1. Exige tratamientos estadísticos muy complejos 2. Hay pérdida de precisión 3. Pérdida del carácter aleatorio del muestreo	Mínimas
	Muestras no probabilísticas (No se conocen las probabilidades de cada individuo o elemento de ser incluidos en la muestra)	Casual	Entrevistar a los individuos, hasta un cierto número, de forma casual (por ejemplo, los que pasan por una esquina)		1. Exige personal menos entrenado y de costo menor	1. Presenta serio obstáculo a la generalización	Mínimas
Intencional		Seleccionar casos típicos del universo según el criterio de un experto		1. Idem	1. Idem	Mínimas	
Cuotas		Cada entrevistador debe entrevistar una cierta cuota de individuos de cada categoría (por ejemplo hombres y mujeres)		1. Idem	1. Idem	Mínimas	
Muestras para probar hipótesis estadísticas	Calcular la matriz de datos más compleja a analizar. Utilizar la fórmula: $M = r_1 \cdot r_2 \cdot r_3 \dots r_k$ Proceder a llenar las casillas ya sea mediante el sistema de cuotas, o por alguna técnica de muestreo estratificado	Ventajas Sirven para probar hipótesis, tienen en cuenta el análisis y las técnicas analíticas		Desventajas Exigen conocimiento previo tanto de características poblacionales cuanto del tipo de análisis que se va a utilizar. Algunas veces no es posible establecer generalizaciones			

Muestras no predispuestas. Son aquellas cuya probabilidad de extracción es conocida. Hay dos muestras de este tipo: muestras cuya probabilidad de ser extraídas es

cero o uno (finalistas); o muestras cuya probabilidad de ser extraídas es diferente de cero o de uno (probabilísticas).

El cuadro 1 sintetiza los tipos de muestra que vamos a exponer en este capítulo, los procedimientos básicos para su obtención y sus ventajas y desventajas. Las muestras que vamos a considerar son:

A) Muestras probabilísticas: 1) Muestra simple al azar. 2) Muestra sistemática. 3) Muestra probabilística (proporcional y no-proporcional). 4) Muestra por conglomerados.

B) Muestras no-probabilísticas: 1) Muestra casual. 2) Muestra intencional. 3) Muestra por cuotas.

C) Muestras para probar hipótesis sustantivas.

A) MUESTRAS PROBABILÍSTICAS

En este tipo de muestra todos los elementos que componen el conjunto total o universo tienen una probabilidad conocida de ser incluidos en la muestra. Describiremos 4 tipos de muestras probabilísticas: 1) Muestreo simple al azar. 2) Muestreo sistemático. 3) Muestreo estratificado. 4) Muestreo por conglomerados.³

1) El muestreo simple al azar

En el muestreo simple al azar, todas las muestras y todos los elementos tienen la misma probabilidad de ser seleccionados. Además, como mencionamos anteriormente, esa probabilidad es conocida y diferente de cero y de uno.

La probabilidad para un individuo o elemento de ser extraído en la muestra resulta de la siguiente relación:

$$p = \frac{1}{N}$$

Donde:

N = tamaño de la población

Por ejemplo: en el sistema de lotería en México, la probabilidad de cualquier número para ser extraído con el premio mayor es igual a:

$$p = \frac{1}{50\,000} = 0.00002$$

En el caso de la probabilidad de una muestra particular de ser extraída es igual a:

$$p = \frac{1}{\frac{N!}{n!(N-n)!}}$$

Donde:

N = tamaño de la población

n = tamaño de la muestra

$!$ = factorial

Es decir, la ecuación nos va a señalar la probabilidad que la combinación de n elementos tiene de ser seleccionada. Y esa probabilidad es la misma para cada una de las muestras posibles de tamaño n . En el caso de una muestra de tamaño 10, cuya población tiene un tamaño de 100, la probabilidad para cualquier muestra de ser extraída sería igual a:

$$p = \frac{1}{\frac{100!}{10!(100-10)!}} = \frac{1}{124\,603\,388\,140}$$

El muestreo aleatorio simple es el más sencillo de todos los tipos de muestreo y, cuando es de tamaño grande, no solamente resulta una muestra probabilística, sino también una muestra por cuotas. Además, cualquiera que sea el tipo de muestreo probabilístico, en algún nivel de éste hay que recurrir al muestreo aleatorio simple (ver muestra estratificada, por ejemplo).

¿Cómo se decide el tamaño de la muestra?

Dijimos más arriba que a través del muestreo se obtienen reducciones en los costos, es decir que, mientras más pequeña sea la muestra, menores serán los costos de la investigación en términos de esfuerzo, dinero, etc. Ahora bien, el problema siguiente es que, a medida que la muestra es más pequeña, la probabilidad de error es mayor; entonces, las decisiones en la determinación del tamaño de la muestra se plantean de la siguiente manera: compatibilizar la disponibilidad de recursos (que normalmente son escasos) con la precisión deseada en las estimaciones.⁴ En otros términos, seleccionar una muestra de tamaño tal que se logre un máximo de precisión, con un tamaño mínimo de muestra.

Galtung⁵ nos presenta una estrategia excelente que consiste en el “principio de las utilidades decrecientes”, es decir, aquel que resulta del siguiente razonamiento: no interesa tanto el tamaño *exacto* de la muestra como la “ganancia” en términos de nivel de significación que el investigador puede obtener aumentando un número determinado de unidades. Y aquí propone dos maneras de razonar:

La primera surge de los métodos estadísticos utilizados para el análisis y

comprobación de las hipótesis; para la determinación del tamaño de la muestra, se parte de diferencias mínimas.

Es decir, la pregunta que debe ser resuelta en un principio es qué nivel de confiabilidad y de significación desea el investigador. Existen tablas para diferentes coeficientes de correlación que especifican para cada nivel de significación el tamaño de muestra necesario. Pero para ello, como bien señala Galtung, es necesario conocer los valores parámetros. Cuando tales valores parámetros no son conocidos, es posible determinar el tamaño por medio de ciertas diferencias. El ejemplo típico señalado por Galtung es el siguiente:

Se está interesado en la significación de *proporciones*. Cuando la proporción en la población es p , y la proporción en la muestra $\frac{x}{m}$, la significación o nivel de confianza es:

$$\overbrace{\frac{x}{m} - p}^d > z \sqrt{\frac{pq}{m}}$$

Donde: $\frac{x}{m}$ = proporción en la muestra

p = proporción en la población

z = ordenada a la curva normal (en este caso con valor 1, es decir, con un nivel de significación de .32).

Los niveles de significación mayormente utilizados en ciencias sociales son de .05 y .01; para ello la ecuación original sería ahora, en el caso de un nivel de significación del 5%:

$$\frac{x}{m} - p > 1.96 \quad z \sqrt{\frac{pq}{m}}$$

si $p = \frac{1}{2}$ tenemos que:

$$d > 1.96 \sqrt{p(1-p)} \cdot \sqrt{\frac{1}{m}} = 1.96 \cdot \sqrt{0.25} \sqrt{\frac{1}{m}}$$

redondeando el 1.96 a 2.00 tenemos que:

$$d^2 = \frac{1}{m}$$

Esta ecuación nos va a señalar cuánto se gana en valores de significación cuando se aumenta la cantidad de unidades en la muestra.

CUADRO 2.⁶ Valores correspondientes de tamaño de la muestra y diferencia significativa
($z = 2$; $p = 0.5$)

m	1	4	9	16	25	36	49	64	81
d	1	0.5	0.33	0.25	0.20	0.17	0.14	0.12	0.11
m	100	121	144	169	225				
d	0.10	0.09	0.08	0.08	0.07				

Es decir que aumentos sucesivos en el tamaño de la muestra producen disminuciones cada vez menores en el error estándar del estimador. En otras palabras, el decremento del error estándar cuando el tamaño de la muestra crece de 50 a 100, por ejemplo, es mayor que cuando la muestra crece de 100 a 150. Esto quiere decir que, a partir de un tamaño de muestra dado, los decrementos del error de estimación requieren tamaños de muestra mayores.

La segunda manera de razonar, similar a la anterior, se hace a partir de la siguiente pregunta: ¿cuál es la utilidad que se desea en términos de disminución del error, para un aumento de tamaño de la muestra? Supongamos que uno decide que cuando no obtenga más que un aumento significativo de 2% o de 1% al aumentar 100 casos va a parar.

Entonces tenemos:

$$d(m) - d(m + 100) = 0.02 \text{ o}$$

$$d(m) - d(m + 100) = 0.01$$

la ecuación generalizada cuando $p = q$ es:

$$2 \sqrt{p(1-p)} \left[\sqrt{\frac{1}{m}} - \sqrt{\frac{1}{m+100}} \right] = 0.02$$

o si se trata de duplicar la muestra, con niveles del 1%

$$2 \sqrt{p(1-p)} \left[\sqrt{\frac{1}{m}} - \sqrt{\frac{1}{2m}} \right] = 0.01$$

Selección de una muestra aleatoria en poblaciones de tamaño finito

La extracción de una muestra aleatoria simple en una población finita requiere un listado de *todos* los elementos de la población. Es decir, una vez definida la población (supongamos que es la de estudiantes en la Universidad Nacional Autónoma de México o las personas que sufragaron en las elecciones presidenciales en 1970 en Chile, o el total de alumnos que concurren a las escuelas primarias en Costa Rica, etc.) es necesario tener un *listado* de la misma. Uno procede a continuación a enumerar los individuos en

la lista, es decir, se les asigna un número. Por ejemplo, si la población total está compuesta de 45 671 individuos, se procede a numerarlos: 00001, 00002, 00003... 04534, 04535... 45669, 45670, 45671. Es importante que en la enumeración el investigador coloque los ceros a la izquierda cuando corresponden, sobre todo si va a utilizar para la extracción tablas aleatorias; en otras palabras, cada número debe tener la misma cantidad de dígitos que el número total de la población. Estas tablas son un conjunto de números naturales que incluyen el cero, presentados en una forma particular y que son obtenidos por medio de algún procedimiento aleatorio (generalmente una computadora) que asegura que esos números no están dispuestos en ningún orden en particular, esto es, con el conocimiento de cualquier número en la tabla no es posible conocer qué número lo precede o antecede.

Presentamos la selección de una muestra aleatoria en las siguientes etapas: 1) Se enumera, siguiendo los criterios señalados más arriba, a todos los individuos que componen el universo. 2) Se determina el tamaño de la muestra. 3) Se selecciona una página cualquiera de la tabla de números aleatorios (la selección misma puede ser hecha al azar). Puesto que en cada página los números aleatorios aparecen dispuestos en columnas, es necesario considerar tantas columnas como dígitos tenga la población. En el ejemplo que mencionábamos más arriba (población = 45 671 casos) tendremos que considerar 5 columnas. Reproducimos a continuación parte de una página de una tabla de números aleatorios: ⁷

98 08 62 48 26	45 24 02 84 04	44 99 90 88 96	39 09 47 34 07
35 44 13 18 80	33 18 51 62 32	41 94 15 09 49	89 43 54 85 81
88 69 54 19 94	37 54 87 30 43	80 95 10 04 06	96 38 27 07 74
20 15 12 33 87	25 01 62 52 98	94 62 46 11 61	79 75 24 91 40
71 96 12 82 96	69 86 10 25 91	74 85 22 05 39	00 38 75 95 79
18 63 33 25 37	98 14 50 65 71	31 01 02 46 74	05 45 56 14 27
.. .. .			
.. .. .			

En el ejemplo que consideramos N es igual a 45 671; por tanto consideramos cinco columnas. Seleccionamos cualesquiera de las columnas (supongamos que partimos de la columna tres). Tendríamos entonces los números 08624; 82645; 24028; 40444; 99908; etc. Esto significa que el individuo 08624 es la primera observación muestral; puesto que el número 82645 no aparece en la lista, lo saltamos. La segunda observación muestral será el sujeto número 24028, y así sucesivamente hasta agotar la cantidad de casos necesarios para el tamaño de la muestra n . Si algún sujeto aparece repetido, también se reemplaza con otra observación muestral. Una vez que se llega al final de una página sin que se haya terminado de completar los casos necesarios para la muestra, es posible pasar a otra página, o simplemente volver al comienzo, empezando ahora con otra columna, diferente de la inicialmente seleccionada.

Como se ve, este tipo de procedimiento muestral implica una cantidad respetable de trabajo, sobre todo cuando la población es grande.

Una salida alternativa que ahorra tiempo consiste en enumerar tanto las páginas del listado de la población como a los individuos incluidos en cada una de las páginas. En el ejemplo seleccionado por nosotros, supongamos que tenemos 46 páginas (45 de ellas con 1 000 sujetos cada una y una página con 671), en las cuales los sujetos aparecen enumerados desde el 001 al 999 en cada página (obviamente en la última página los sujetos estarán numerados del 001 al 671).

Procedemos entonces a trabajar también ahora con 5 columnas, sólo que ahora las dos primeras columnas nos indicarán el número de la página, y las tres siguientes, el número de los sujetos; por ejemplo, si hubiésemos comenzado con la columna seis de nuestra tabla, tendríamos los siguientes números 24826; 45240; 28404; 44999; 08896, etc.; es decir que seleccionamos el sujeto 826 en la página 24; el 240 en la página 45; el 404 en la página 28; el 999 en la página 44, y así sucesivamente. Lo mismo que en el procedimiento anterior saltamos los números que no correspondan a ninguna página, así como aquellos que aparezcan repetidos. La ventaja de este sistema es que se reconoce mucho más rápidamente qué sujetos aparecen repetidos.

Hay otros procedimientos de selección aparte de tablas de números aleatorios, tales como los que se utilizan para la selección de los números a ser premiados en las loterías (en donde existen dos series independientes de selección: una para los números y otra para los premios); pero en este caso hay que recordar que es necesario reponer los números nuevamente en el bolillero, de manera tal que la probabilidad para cualquier bolilla sigue siendo la misma. Esto no se realiza, por ejemplo, en la lotería, es decir, una vez que un número ha sido extraído no es repuesto.

Ventajas y desventajas del muestreo aleatorio simple

Ventajas. a) No supone el conocimiento previo de *ninguna* de las características de la población de la cual se va a extraer la muestra. Esto es, a diferencia del muestreo estratificado, por ejemplo, no es necesario conocer la frecuencia relativa con que se dan las características poblacionales en cada uno de los estratos. Esto significa que una muestra aleatoria simple está libre de los sesgos que se pueden introducir por el uso de ponderaciones incorrectas en las unidades muestrales (ver muestreo estratificado). b) Es relativamente simple determinar la precisión de las estimaciones que se hacen a partir de las observaciones muestrales porque los errores estándar de los estimadores siguen distribuciones de probabilidad conocidas. Esta ventaja resulta del hecho de que la teoría del muestreo aleatorio simple está más desarrollada que ninguna otra. c) Tiende a reflejar todas las características del universo, esto es, cuando el tamaño de la muestra crece se hace cada vez más representativa del universo o población.

Desventajas. a) Supone un listado completo de todas las unidades que componen la población. Obviamente, en muchos casos no se cuenta con una lista completa y actualizada de la población, lo que impide el empleo de este diseño muestral. b) Aun cuando se cuente con este listado, su numeración demanda mucho tiempo y trabajo que pueden ahorrarse si se emplea un diseño muestral distinto. c) Supone un tamaño de

muestra mayor que otros diseños para obtener un mismo nivel de confiabilidad. Esto significa que, para un mismo tamaño de muestra, las estimaciones hechas a partir de una muestra estratificada son más precisas que las mismas estimaciones hechas a partir de una muestra aleatoria simple. *d)* Es probable que las unidades muestreadas (si son individuos que viven en una ciudad, por ejemplo) queden muy distantes unas de otras, con lo cual el costo para obtener la información de estas unidades crece con la dispersión espacial de las mismas. Otros diseños muestrales permiten reducir al mínimo esta dispersión.

Ejemplo de una muestra aleatoria simple

A continuación se ilustran con un ejemplo las etapas ya señaladas en la extracción de una muestra aleatoria simple. La población está constituida por 5015 individuos que componen las 1003 mesas (de cinco vocales cada una) correspondientes a las doce comunas de Santiago de Chile. En cada una de estas comunas hay un conjunto de mesas constituidas por mujeres, y otro constituido por varones, además de las mesas formadas por extranjeros.

De esta población se extrajeron dos muestras aleatorias simples de 200 individuos cada una. La primera de estas muestras (200 a) se extrajo directamente, vale decir, sin emplear el método de enumeración de páginas descrito en la página 99. La segunda se extrajo después de haber enumerado las páginas del listado de la población (200 b). En este caso, se formaron seis páginas con alrededor de 840 individuos cada una.

Por medio de las dos muestras se estimaron dos parámetros: la proporción de varones en la población (p) y el promedio aritmético (μ) de los puntajes correspondientes a una variable cuyo recorrido era 0, 1, 2, ... 9. Esto se hizo con el objeto de comparar las estimaciones obtenidas a partir de los dos métodos de extracción de las observaciones. La información contenida en cada una de las muestras fue la siguiente:

CUADRO 2

	<i>Muestra 200 a</i>	<i>Muestra 200 b</i>
Número de hombres en la muestra	95	98
Proporción de hombres en la muestra	0.475	0.490
Error estándar de proporción muestral	0.0354	0.0361
Media aritmética de los puntajes en la muestra	5.20	5.12
Error estándar de la media aritmética	0.200	0.195

Los intervalos de confianza de 95% para (p) y (μ) fueron los siguientes:

	<i>Muestra 200 a</i>	<i>Muestra 200 b</i>
para (p):	$0.405 < p < 0.545$	$0.420 < p < 0.560$
para (μ):	$4.81 < \mu < 5.19$	$4.78 < \mu < 5.46$

2) Muestreo sistemático

Es muy similar al muestreo aleatorio simple. Hay que confeccionar un listado de todos los elementos que incluye la población; una vez en posesión del listado de la población, la diferencia estriba en el método para la selección de los casos. En la muestra aleatoria simple utilizábamos tablas aleatorias o procedimientos similares; aquí la selección se realiza por un procedimiento más mecánico que representa un gran ahorro de tiempo.

Supongamos que la población es de 50 mil casos. Una vez hecho el listado y determinado el tamaño de la muestra (por ejemplo, de 1 000 casos) se procede de la siguiente manera: *a)* se selecciona al azar un número comprendido en la cantidad que resulte de dividir el tamaño de la población entre el tamaño de la muestra (en nuestro caso, entre 01 y 50), y a partir de ese número se seleccionan las unidades de la siguiente manera: *b)* supongamos que el número seleccionado es el número 15; entonces extraigo los casos 15, 65, 115, 165, 215... etc. (es decir, cada 50 casos) hasta completar las 1 000 unidades.

Ejemplo: Si se utiliza el directorio de teléfonos o cualquier listado similar, se enumera a todos los sujetos listados; se selecciona una primera fracción que incluya todas las páginas con personas listadas, y a partir de la primera página se extraen sujetos con un número fijo (cada 20, cada 87, cada 140, o cualquier otro), hasta completar la cantidad de casos.

Ventajas y desventajas

Ventajas. *a)* Las ventajas técnicas son similares a las del muestreo aleatorio simple. *b)* Tiene una gran ventaja económica ya que facilita la selección de la muestra, sobre todo en aquellos casos en los que ya existe un listado.

Desventajas. *a)* Desde el punto de vista estrictamente estadístico, este tipo de muestreo no es probabilístico ya que, si bien es correcto que, una vez iniciado el muestreo las probabilidades son las mismas para la primera selección —es decir, estrictamente para los n sujetos que están incluidos en el número fijo—, una vez elegido este número, la muestra pasa a ser finalista (en nuestro ejemplo, la probabilidad de los sujetos 65, 115, 165, etc., es 1, mientras que la probabilidad de cualquier otro sujeto en el listado es 0). *b)* La desventaja más sobresaliente es la que puede resultar de los sesgos propios del listado que estamos utilizando. El directorio de teléfonos, por ejemplo, dispone de espacios mayores para el listado de comercios, empresas, etc., que para individuos particulares. De manera que, si no tomamos en cuenta esto, introducimos en nuestra muestra el sesgo del listado.

3) Muestreo estratificado

Este tipo de muestra es conveniente cuando la población o universo puede dividirse en categorías, estratos o grupos que tienen un interés analítico, y que por razones teóricas y empíricas presentan diferencias entre ellos. La ventaja que ofrece la estratificación es que permite una mayor homogeneización de la muestra final.

Por ejemplo, se puede estratificar una población según sexo, edad, *status* socioeconómico, nivel ocupacional, características de personalidad, étnicas, educacionales, etc. O bien, si las unidades de análisis son comunidades, se puede estratificarlas según su nivel de desarrollo socioeconómico, características de producción, geográficas, etc. Es decir que la estratificación en sí, la definición de cada uno de los estratos, es un problema de propósitos de investigación, tipo de preguntas que se quiere responder y teoría sustantiva.

Una vez definidos los estratos y dividido el universo en estratos o valores según una o más variables, es posible diferenciar *dos tipos de muestras estratificadas*: a) Muestra estratificada proporcional, y b) Muestra estratificada no proporcional.

La muestra estratificada proporcional es aquella en la cual la fracción de muestreo es igual para cada estrato; si existen diferencias en las fracciones de muestreo, se llama *no proporcional*.

Una vez determinada la proporción, se seleccionan las muestras *dentro de cada estrato* según los procedimientos del muestreo aleatorio simple señalado más arriba.

Ejemplo 1: Muestra estratificada proporcional

Supongamos que el propósito de la investigación es probar algunas hipótesis acerca del rendimiento escolar. Nuestro sistema de hipótesis plantea que existen diferencias significativas en cuanto al rendimiento escolar de los niños, según su extracción de clase, y una serie de otros factores que aquí vamos a dejar de lado para simplificar el ejemplo. Definido nuestro interés en estratificar a la población escolar según su clase social, supongamos ahora que estamos en una escuela concreta a la cual concurren 500 niños; determinamos la composición de clase de la escuela y el resultado es el siguiente: 50 niños provienen de clase alta; 300 niños provienen de clase media, y 150 niños provienen de clase baja.

Supongamos que el tamaño de muestra a seleccionar es del 10% del universo, esto es, 50 casos. Si utilizamos muestreo aleatorio simple, la probabilidad de que la muestra resultante contenga exactamente un 10% de cada estrato (30 niños de clase media, 15 de clase baja y 5 de clase alta) es muy baja. Nosotros queremos garantizar que vamos a obtener exactamente esa cantidad. Entonces, en este caso vamos a trabajar con un muestreo aleatorio proporcional, es decir, la fracción de muestreo para cada estrato será exactamente la misma: 10%. El método de selección entonces es el siguiente:

a) Estratifico a la población en tres clases: niños de clase media, alta y baja, respectivamente.

b) Confecciono un listado independiente para cada estrato, enumerando a mis sujetos

(en la lista *A* —de clase media— unos niños estarán enumerados del 001 al 300; en la lista *B*, de clase baja, estarán enumerados del 001 al 150; en la lista *C* estarán enumerados del 01 al 50).

c) Procedo a la selección, por los métodos indicados en el muestreo aleatorio simple, de 30 casos en la lista *A*, 15 casos en la lista *B*, y 5 casos en la lista *C*. Con ello logro que la muestra de 50 casos sea una réplica proporcional exacta del universo, en lo que a clase social se refiere. Para los propósitos de generalización las muestras estratificadas proporcionales no ofrecen complicaciones en el cálculo.

Ejemplo 2: Muestra estratificada no proporcional

Estamos en la misma escuela del Ejemplo 1. Estamos interesados en comparar los rendimientos de las distintas clases. Ocurre que con la muestra que obtuvimos mediante el muestreo aleatorio estratificado proporcional no contamos con suficientes casos de niños de extracción de clase alta, ya que tenemos solamente 5.

Decido entonces que para los propósitos de mi análisis voy a necesitar 25 casos en cada estrato, es decir que mi muestra va a tener un total de 75 casos. Tendré entonces un *muestreo estratificado no proporcional*, ya que no voy a respetar la proporción original en el universo.

CUADRO 3

<i>Universo</i>		<i>Muestra</i>	
<i>Clases</i>	<i>Sujetos</i>	<i>Sujetos</i>	<i>Fracción de muestreo</i>
Clase alta	50	25	50%
Clase media	300	25	8.3%
Clase baja	150	25	16.7%
	<u>500</u>	<u>75</u>	

Una vez decidida la cantidad total de la muestra y la cantidad de casos en cada estrato, procedo a seleccionarlos según los mismos procedimientos indicados para el muestreo aleatorio simple, tomando 25 casos de cada lista (*A*, *B* y *C*). Obtengo entonces el cuadro 3.

Esta tabla es importante para los propósitos de generalización a la población, ya que ahora tengo que *ponderar* las diferentes fracciones de muestreo.

Para ello, en el caso de la media aritmética, por ejemplo, hay que operar de la siguiente forma: a) Calcular la media aritmética para cada estrato; b) Ponderarlas según el tamaño relativo del estrato.

CUADRO 4

	<i>Estratos</i>			
	<i>Alto</i>	<i>Medio</i>	<i>Bajo</i>	<i>Total</i>
Tamaño del estrato	50	300	150	500
<i>Peso</i> (P_i)	0.10	0.60	0.30	1.00
Tamaño de la muestra	25	25	25	
Media aritmética (M_i)	24	20	17	
Desviación estándar (σ_i)	5	4	7	

El cálculo de la media muestral será realizado de acuerdo a la siguiente fórmula:

$$M_T = \sum P_i M_i$$

Donde: M_i = media aritmética de cada estrato

P_i = peso de cada estrato

En nuestro caso:

$$M_T = (20 \times 0.60) + (17 \times 0.30) + (24 \times 0.10) = 19.5$$

En el caso del error estándar:

$$\sigma_T = (4 \times 0.60) + (7 \times 0.30) + (5 \times 0.10) = 5.0$$

Ventajas y desventajas

Ventajas (de las ofrecidas por el muestreo aleatorio simple). *a)* La muestra es más homogénea, garantizando la representatividad. *b)* Elimina los errores en la estimación que son producto de diferencias entre estratos. *c)* Una ventaja adicional ofrecida por el muestreo estratificado no proporcional es la de posibilitar un mejor conocimiento de grupos pequeños (en relación a la cantidad total de casos).

Desventajas. *a)* Supone el conocimiento *previo* de las características de la población, a partir de las cuales se estratifica. *b)* Son de costo más elevado que las aleatorias simples, en dinero y energía. *c)* Exige tratamientos estadísticos de cálculo más complejo. *d)* Pueden existir dificultades en la determinación de estratos homogéneos. *e)* La muestra estratificada proporcional algunas veces puede no proveer un número suficiente de casos para análisis comparativos inter-estratos.

4) Muestreo por conglomerados

Muchas investigaciones en ciencias sociales tienen como objeto de estudio unidades tales como naciones, estados y similares, que admiten subdivisiones o que ya contienen distintos conglomerados.

En este sentido hay bastante similitud entre las muestras estratificadas y las muestras por conglomerados, aunque existen diferencias importantes en cuanto a los métodos de selección en uno y otro caso.

En términos generales, el investigador considera el muestreo por conglomerados en aquellos casos en los cuales la población a estudiar está dispersa a lo largo de áreas geográficas extensas o situaciones similares, donde los costos para alcanzar las unidades resultan muy elevados.

Los procedimientos para la selección de la muestra en este tipo de muestreo son los siguientes:

a) Es necesario dividir a la población en conglomerados lo más homogéneos posible. Esta división por conglomerados puede hacerse en varios niveles; es decir, se puede operar en varios escalones o etapas.

b) Una vez determinados los conglomerados, se selecciona al azar del primer nivel de conglomerados una proporción determinada de los mismos (si seleccionamos todas las unidades dentro de estos conglomerados, hablamos de *muestra* de escalón o de una sola etapa).

c) Con los conglomerados seleccionados en *b*, procedemos a una nueva selección de conglomerados al interior de cada uno de ellos (nuevamente, si tomamos todos los casos dentro de este segundo nivel, hablamos de *muestras de 2 escalones*.)

d) Si procedemos por el mismo método a seleccionar en un tercer, o cuarto nivel, hablamos de *muestras de escalones múltiples*.

Téngase bien claro que, una vez determinados los conglomerados, se procede a una selección aleatoria *entre* conglomerados, y una vez seleccionados éstos, *tomamos todas las unidades de cada conglomerado seleccionado*. Ésta es la diferencia entre una muestra por conglomerados y una muestra estratificada. En la muestra estratificada, una vez determinados los estratos, seleccionábamos los casos de cada estrato (proporcional o no proporcionalmente).

Un ejemplo nos ayudará mejor a seguir con mayor detalle el método para la selección de unidades en este tipo de muestreo por conglomerados. Y vamos a elegir un ejemplo particularmente extenso para hacer más clara la exposición.

Supongamos que nos interesa determinar ciertas características sociopsicológicas en los Estados Unidos Mexicanos, es decir que la unidad de análisis es un país. México tiene una extensión de unos 2 millones de kilómetros cuadrados y aproximadamente 60 millones de habitantes [datos de 1979. E.]. Si para los cálculos vamos a partir de datos agregados sobre características de los mexicanos, plantear una muestra aleatoria simple o aun una muestra estratificada implicaría el problema de la dispersión de las unidades a lo largo del país. Por lo tanto, nos decidimos por una muestra por conglomerados. Y vamos a realizarla en varios *escalones*:

i) *En el primer escalón*, vamos a tomar la división política de México en 32 Estados.⁸

Tendríamos entonces el mapa que aparece en la página siguiente.

ii) Obsérvese que numeramos los 32 estados, conglomerados ahora. Procedo a continuación a seleccionar al azar 11 de estos conglomerados (aproximadamente un tercio). El sistema de selección puede ser puramente al azar; es decir, utilizamos una tabla, o simplemente colocamos 32 bolitas, y vamos extrayendo de una en una hasta completar 11 (no olvidar de reponer); o como hemos elegido en este caso, seguir el sistema de las “agujas de reloj”, es decir, seleccionamos de entre las tres primeras al azar una y luego procedemos como en el caso del muestreo sistemático a seleccionar cada tercera. En otras palabras, selecciono al azar la unidad 1, 2 y 3; supongamos que resulta elegida la número 3, entonces saldrán automáticamente seleccionadas las unidades 3, 6, 9, 12, 15, 18, 21, 24, 27, 30 y 1. Con este sistema busco tener representados todos los segmentos geográficos a lo largo de la República. Éstos son los estados o conglomerados que aparecen sombreados en el mapa. Si se tratara de una muestra de un *escalón*, censaría a todos los habitantes de los estados seleccionados; pero vamos a proceder a nuevos *escalones*, y para representarlo vamos a utilizar el Distrito Federal, que suponemos que ha salido seleccionado.

iii) Vamos a utilizar nuevamente aquí las 25 divisiones zonales (13 municipios y 12 cuarteles) utilizadas en el Censo Nacional de Población de 1970 (ver mapa 2).

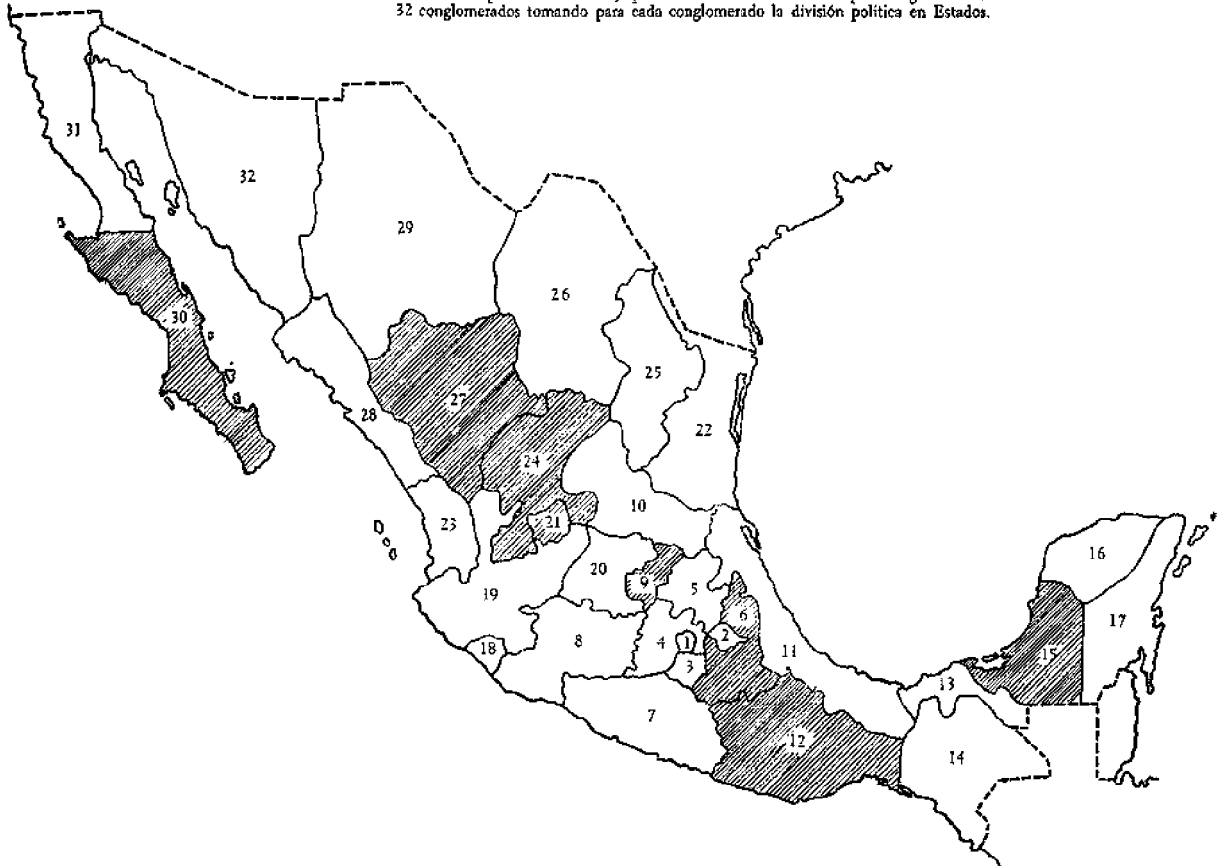
El Censo Nacional de Población de 1970, del cual tomamos la división del Distrito Federal en 13 delegaciones censales, basó la jurisdicción en criterios tales como superficie en kilómetros cuadrados, número de localidades, población estimada para 1970, vías de comunicación existentes y accidentes topográficos.

La Zona 1, que concentra la mayor cantidad de habitantes (ver cuadro 5), fue dividida en cuarteles (12 en total).

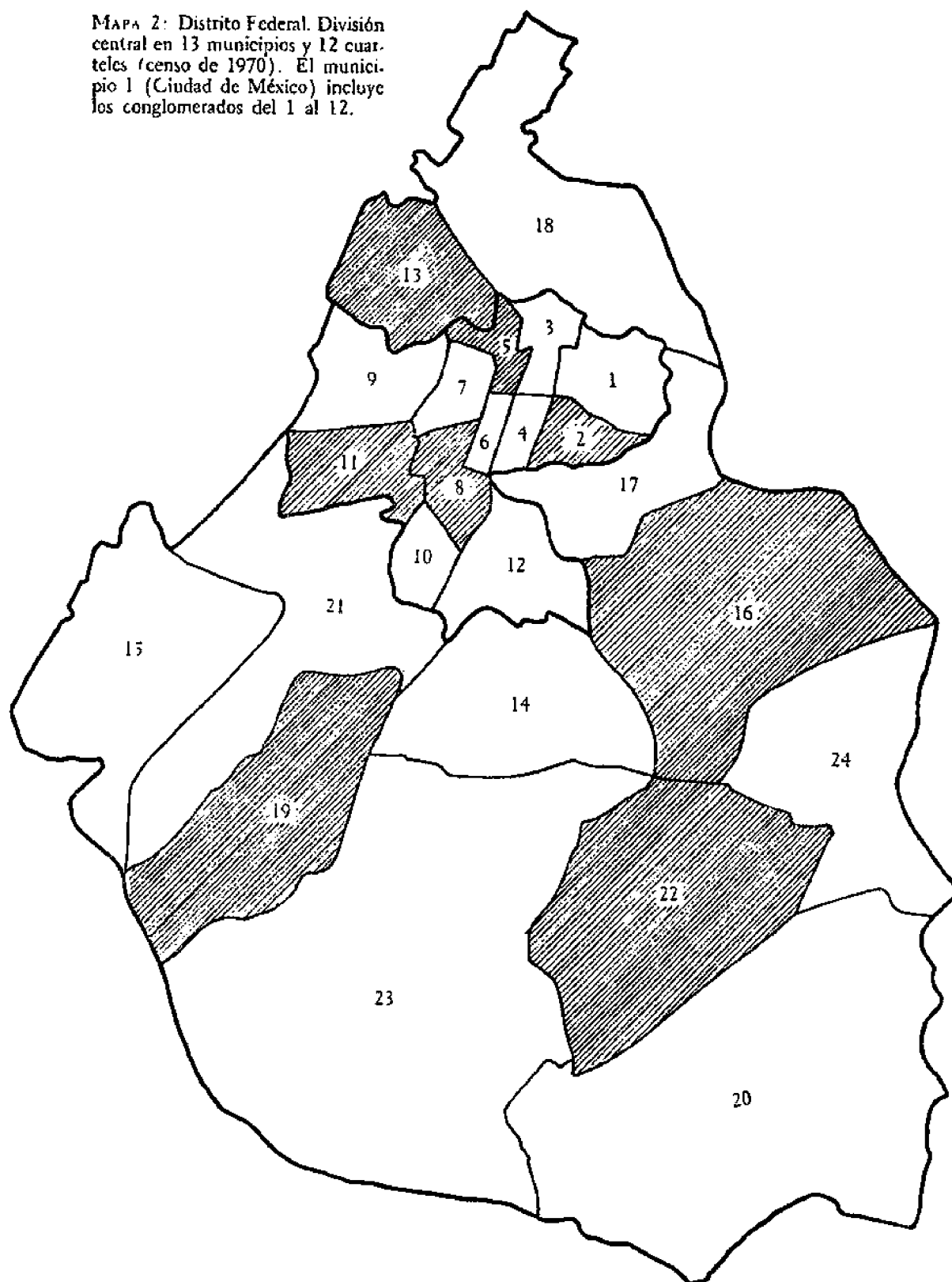
Para los fines del muestreo por conglomerado, cada cuartel de la primera división censal 1 (Ciudad de México) será tratado como un conglomerado aparte. En este caso habría que subdividir el mapa en 24 conglomerados.

iv) Seleccionaremos nuevamente por el sistema de las “agujas de reloj” un tercio de las zonas. Las zonas sombreadas corresponden a la selección que se produjo por efecto de la extracción del número 2, en la primera selección al azar entre las tres primeras zonas. Recordemos que este procedimiento se repite en el resto de los 10 conglomerados seleccionados en el primer *escalón*. Si nos detenemos aquí y censamos a todos los individuos incluidos en las zonas resultantes, tenemos una muestra de dos *escalones*. Pero vamos a seguir más adelante. Para ello continuamos utilizando el Distrito Federal, aunque ahora nos concentremos en la Zona 2.

MAPA 1: República Mexicana; primer escalón en muestra por conglomerados.
32 conglomerados tomando para cada conglomerado la división política en Estados.



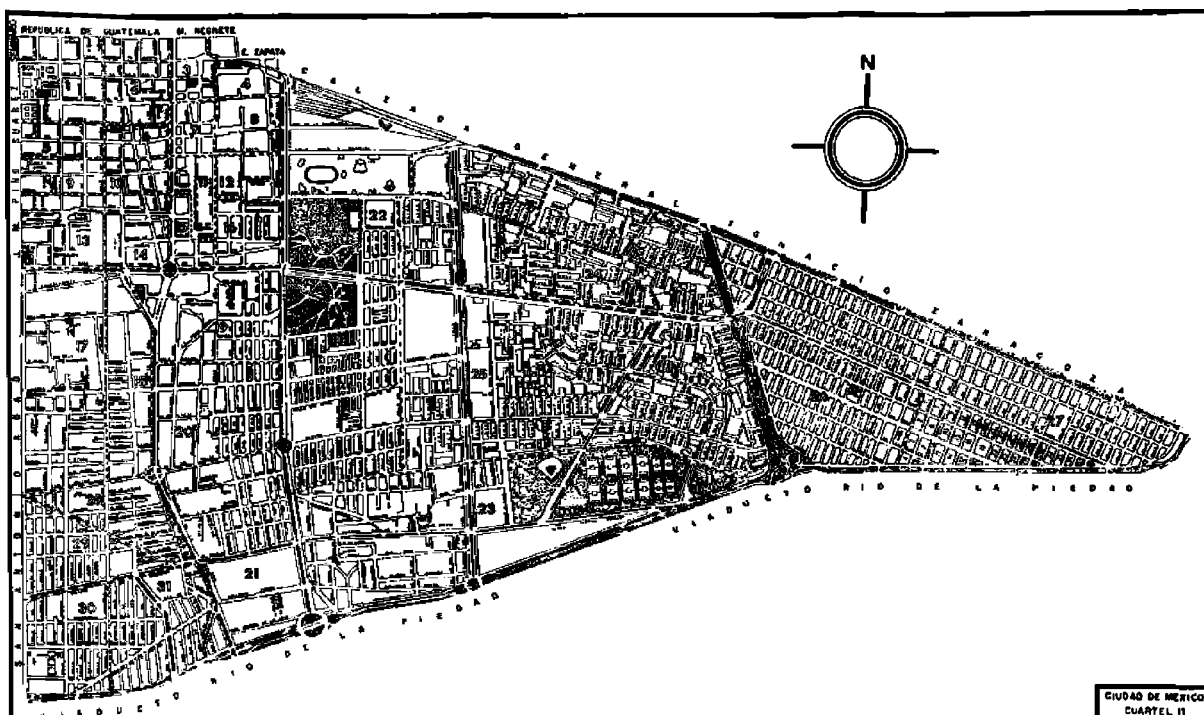
MAPA 2: Distrito Federal. División central en 13 municipios y 12 cuarteles (censo de 1970). El municipio 1 (Ciudad de México) incluye los conglomerados del 1 al 12.



CUADRO 5. *Divisiones censales y cuarteles en el Distrito Federal para el Censo Nacional de Población de 1970*

<i>División censal</i>	<i>Habitantes</i>
1. Ciudad de México	2 902 969
<i>Cuartel:</i>	
I. 584 879	
II. 306 530	
III. 141 347	
IV. 104 156	
V. 112 779	
VI. 97 675	
VII. 166 577	
VIII. 231 016	
IX. 429 664	
X. 199 653	
XI. 226 983	
XII. 301 710	
2. Azcapotzalco	534 554
3. Coyoacán	339 446
4. Guajimalpa	36 200
5. Gustavo A. Madero	1 186 107
6. Iztacalco	477 331
7. Iztapalapa	522 095
8. Magdalena Contreras	75 429
9. Milpa Alta	33 694
10. Obregón	456 709
11. Tláhuac	62 419
12. Tlalpan	130 719
13. Xochimilco	116 493
<i>Total</i>	6 874 165

MAPA 3



v) En el tercer *escalón*, vamos a definir conglomerados en función de la cantidad de habitantes. Para ello vamos a dividir la zona en x áreas de aproximadamente 5 mil habitantes cada una. Las divisiones que figuran a continuación en el mapa 3 no tienen que ver con los 5 mil habitantes, sino que se hace con el propósito de ilustración simplemente.

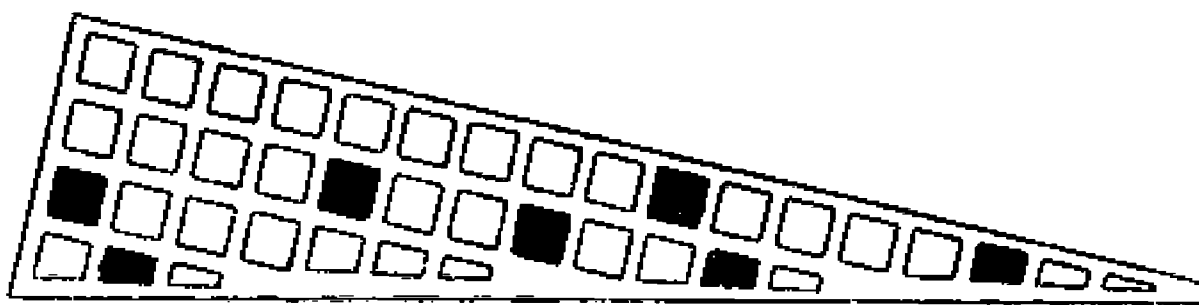
vi) Procedo nuevamente a seleccionar, por los mismos métodos señalados en cada *escalón* anterior, digamos ahora a la cuarta parte de las subzonas o conglomerados de tercer nivel. Las áreas sombreadas corresponden a los conglomerados seleccionados. Procedo en forma idéntica con el resto de las 7 zonas seleccionadas en el segundo nivel. Si el investigador decide detenerse en este *escalón*, debería tomar entonces todos los casos dentro de cada zona. Pero nosotros vamos a proceder más adelante.

vii) Tomamos, para ejemplificar, el conglomerado número 32. Voy a seleccionar ahora manzanas dentro del conglomerado. Para ello las enumero y procedo en forma idéntica a los casos anteriores. El criterio para seleccionar las manzanas no seguirá el sistema de las agujas de reloj, sino que seleccionaremos al azar 20% de ellas. Nuevamente, de detenerme aquí censaría a todos los habitantes dentro de las manzanas seleccionadas. Podría seguir más adelante y seleccionar ahora edificios y, una vez seleccionados éstos, censar a todos. El mapa 4 señala las manzanas seleccionadas en el cuarto *escalón*.

Como el lector podrá anticipar, el sistema de cálculo para las estimaciones de parámetros es complicado, ya que hay que tomar en cuenta los factores de proporcionalidad en cada *escalón*.

MAPA. 4. Ejemplo del 4º *escalón*. Corresponde al conglomerado 32 del cuartel 2 de la

Ciudad de México. El sombreado corresponde a las manzanas seleccionadas.



Ventajas y desventajas

Ventajas. a) La ventaja principal de este tipo de muestreo es la concentración de las unidades en áreas, lo que evita el desplazamiento a lo largo de áreas muy extensas, representando consecuentemente una economía de gastos en términos de energías, costos y tiempo. b) Otras ventajas son similares a las ofrecidas por los tipos de muestras indicadas más arriba.

Desventajas. a) Exige tratamientos estadísticos algo complejos. b) En comparación con el muestreo aleatorio simple o el estratificado, representa pérdidas en precisión. c) La muestra pierde carácter de probabilística, en la medida en que, una vez seleccionados los conglomerados primeros, la probabilidad de ser extraídos que tienen los individuos del resto de los conglomerados es cero.

En síntesis

Los cuatro tipos de muestras aleatorias que presentamos pueden ser aplicados en situaciones diferentes, siendo también posible hacer combinaciones entre ellos. Es decir, se puede partir de muestras de conglomerados e incluir en los escalones finales criterios de muestra estratificada, por ejemplo, aunque por supuesto esto implica complicaciones adicionales en términos de cálculos estadísticos en el momento de la estimación de los parámetros.

La muestra aleatoria simple, cuando el número de casos es abundante, en términos generales debe ser preferida al resto de las otras muestras, particularmente cuando no se conoce lo suficiente sobre los parámetros poblacionales. Cuando es posible estratificar hay que hacerlo, ya que esto homogeneiza la muestra y los cálculos consiguientes resultan de mayor riqueza y representatividad.

La muestra por conglomerados, si bien tiene algunas limitaciones en términos de inferencia, representa muchas veces la única salida cuando los recursos disponibles no son suficientes para plantear cualquiera de las dos primeras alternativas.

La muestra sistemática debe ser utilizada únicamente en aquellos casos en que resulte imposible o demasiado costoso aplicar una muestra aleatoria simple, o como una etapa secundaria de un muestreo complejo.

Describimos ahora rápidamente algunos tipos de muestras no probabilísticas, concentrándonos principalmente en las “muestras” para probar hipótesis sustantivas (ver Galtung),⁹ por su interés para la investigación.

B) MUESTRAS NO PROBABILÍSTICAS

Consideramos aquí 3 tipos: 1) Muestras casuales. 2) Muestras intencionales. 3) Muestras por cuotas.

Recordemos que estas muestras tienen poco valor en términos de estimación, ya que no es posible a partir de ellas calcular el error de estimación de parámetros. Sin embargo, muchas veces resultan de utilidad en términos de inmersión por parte del investigador en el tema. Un caso especial es representado por las muestras para proponer hipótesis sustantivas, cuyo tratamiento en detalle se hará en la sección correspondiente.

1) Las muestras casuales

Son la técnica favorita de muestreo de reporteros de canales de televisión y de la radio, amén de algunas agencias de investigación de mercado. La técnica consiste simplemente en entrevistar a sujetos en forma casual —por ejemplo, uno de cada 10 individuos que pasan por una esquina determinada de la ciudad—. En su grado más simple, este tipo de muestreo y las investigaciones en las cuales se aplica requieren personal menos calificado, de manera que el costo resulta muy reducido. Sin embargo, como señalamos varias veces, a partir de ellas es casi imposible hacer alguna generalización válida que vaya más allá de las características de los individuos que andan a pie, en determinadas horas, por determinadas calles, etcétera.

2) Las muestras intencionales

Son el producto de una selección de casos según el criterio de algún experto; de esta forma se seleccionan algunos casos que resultan “típicos”. De nuevo, estas muestras no son totalmente inútiles, porque resultan de importancia en las etapas exploratorias de la investigación, sobre todo si se utilizan estos casos como “informantes clave” sobre situaciones específicas (ver diagrama en el capítulo I, “La organización de un *survey*”).

3) Las muestras por cuotas

Son en alguna medida una especie de muestras, estratificadas, y son muy utilizadas por algunas agencias de investigación de mercado. A una serie de entrevistadores (profesionales o no) se les fija una cuota de individuos a entrevistar, especificándoles sus características (por ejemplo: varones, casados, propietarios de automóvil, no mayores de 50 años ni menores de 25). Cada entrevistador selecciona por su cuenta y entrevista a los sujetos según un cuestionario, hasta completar su cuota. Como debe ser evidente a estas alturas, el sesgo del entrevistador es una de las desventajas potenciales más importantes

que tiene este tipo de muestreo.

C) MUESTRAS PARA PROBAR HIPÓTESIS SUSTANTIVAS

Habíamos señalado en el comienzo del capítulo que no toda muestra tiene como propósito la estimación de parámetros poblacionales a partir de las características de la muestra, sino que las muestras podrían tener también interés analítico por sí mismas. Siguiendo de muy cerca a Galtung, vamos a dedicar algún espacio a este tipo de muestra, ya que por su uso es una de las más difundidas.

Las muestras para poner a prueba hipótesis sustantivas deben seleccionarse de manera tal que contengan el tipo de elementos a los cuales hacen referencia las proposiciones en que está interesado el investigador. Es decir, en este caso el investigador no está tan interesado en la generalización como en la relación específica entre variables; de manera que quiere garantizar que su muestra contenga unidades suficientes de un tipo determinado. Es decir que la muestra debe ser suficientemente heterogénea. Galtung es muy claro en diferenciar las muestras casuales e intencionales y por cuota, señalando los problemas a que cada una de ellas puede dar lugar. Aquí se trataría de una variación de la muestra por cuota, aunque ahora teniendo cuidado especial en la selección de las unidades para cada casillero.

En general, dice Galtung, el investigador debe responder a tres preguntas para decidir el tamaño de su muestra: *a)* ¿cuántas variables quiere el investigador analizar simultáneamente?; *b)* ¿cuál es el número máximo de valores que desea utilizar por variable?, y *c)* dadas las técnicas analíticas a utilizar, ¿cuál es el valor mínimo por celda que necesita?

Es decir que por medio de la respuesta a las dos primeras preguntas va a determinar el tamaño de la matriz de datos (espacio de atributos, en lenguaje de Barton), y por la última va a satisfacer los requisitos referentes a la prueba de hipótesis estadísticas.

El procedimiento para la determinación del tamaño de la matriz de datos depende entonces de la cantidad de variables (n) y del número de valores en cada variable (r). La fórmula para la determinación del tamaño de la matriz será:

$$\text{Matriz} = r_1 \cdot r_2 \cdot r_3 \cdot \dots \cdot r_n$$

El siguiente cuadro especifica la cantidad de celdas que resultan de combinaciones para diferentes números de variables, con valores iguales para cada una de ellas. Los números que figuran entre paréntesis corresponden a 10 y 20 casos por celda, respectivamente.

CUADRO 6. Tamaño de la matriz para diferentes combinaciones de variables con idéntica cantidad de alternativas

Número de variables (n)	Número de valores en las variables (r)			
	2	3	4	5
1	2 (20-40)	3 (30-60)	4 (40-80)	5 (50-100)
2	4 (40-80)	9 (90-180)	16 (160-320)	25 (250-500)
3	8 (80-160)	27 (270-540)	64 (640-1 280)	125 (1 250-2 500)
4	16 (160-320)	81 (810-1 620)	256 (2 560-5 110)	625 (6 520-12 500)
.				
.				

Tomemos un ejemplo para la determinación del tamaño de la matriz cuando los valores de algunas variables son diferentes. Supongamos que la combinación más compleja de variables en la investigación va a ser el resultado de un análisis de χ^2 o de proporciones, para la prueba de la hipótesis, de las siguientes variables combinadas: *a) Sexo*: que tendrá dos valores: masculino-femenino. *b) Educación*: que tendrá tres valores: alta-media-baja. *c) Religiosidad*: que tendrá tres valores: alta-media-baja. *d) Participación política*: que tendrá 4 valores: alta-media alta-media baja-baja.

La matriz de datos para el análisis simultáneo de las cuatro variables tendrá entonces 72 celdas ($2 \times 3 \times 3 \times 4$), para lo cual necesitaremos al menos una muestra de 720 casos (10 casos por celda). La matriz final tendrá la siguiente forma:

CUADRO 7. Matriz de datos resultante para la combinación de las variables. Sexo (masculino-femenino); Educación: (alta-media-baja); Religiosidad: (alta-media-baja); Participación Política: alta-media alta-media baja-baja)

	Varones									Mujeres								
	Educación									Educación								
	Alta			Media			Baja			Alta			Media			Baja		
	Religio.			Religio.			Religio.			Religio.			Religio.			Religio.		
	A	M	B	A	M	B	A	M	B	A	M	B	A	M	B	A	M	B
Participación política																		
Alta																		
Media alta																		
Media baja																		
Baja																		

Es decir que para probar mi hipótesis tengo que tener varones y mujeres, con distintos niveles educacionales y diferentes grados de religiosidad, para medir su influencia sobre la variable participación política.

Galtung presenta una doble línea de argumentación para justificar la cantidad de 10 a 20 casos por casillero. Por la primera demuestra que el número de casos debe ser suficiente para aproximar la distribución muestral a la población. Para ello parte de intervalos de confianza. Por la segunda el argumento está dirigido hacia la necesidad de que los porcentajes que se computen estén sujetos a grandes variaciones, cuando se dan pequeños cambios en las cifras absolutas. Si el cambio en una unidad produce una variación de 5 puntos de porcentaje y ése es el valor máximo que queremos aceptar, entonces esto nos dará una base de 20 cuando hay dos celdas, de 30 cuando hay tres, etcétera.

Para finalizar, queremos insistir en lo siguiente: las muestras para probar hipótesis sustantivas son de mucha utilidad y es posible compatibilizar las exigencias de un muestreo probabilístico con las exigencias para probar hipótesis sustantivas. Las ventajas de plantear el problema del muestreo de esta manera están vinculadas al proceso total de la investigación, ya que obligan al investigador a explicar como hipótesis, y a pensar desde un comienzo en los métodos a utilizar en el análisis.

IV. EL CUESTIONARIO

JORGE PADUA
INGVAR AHMAN

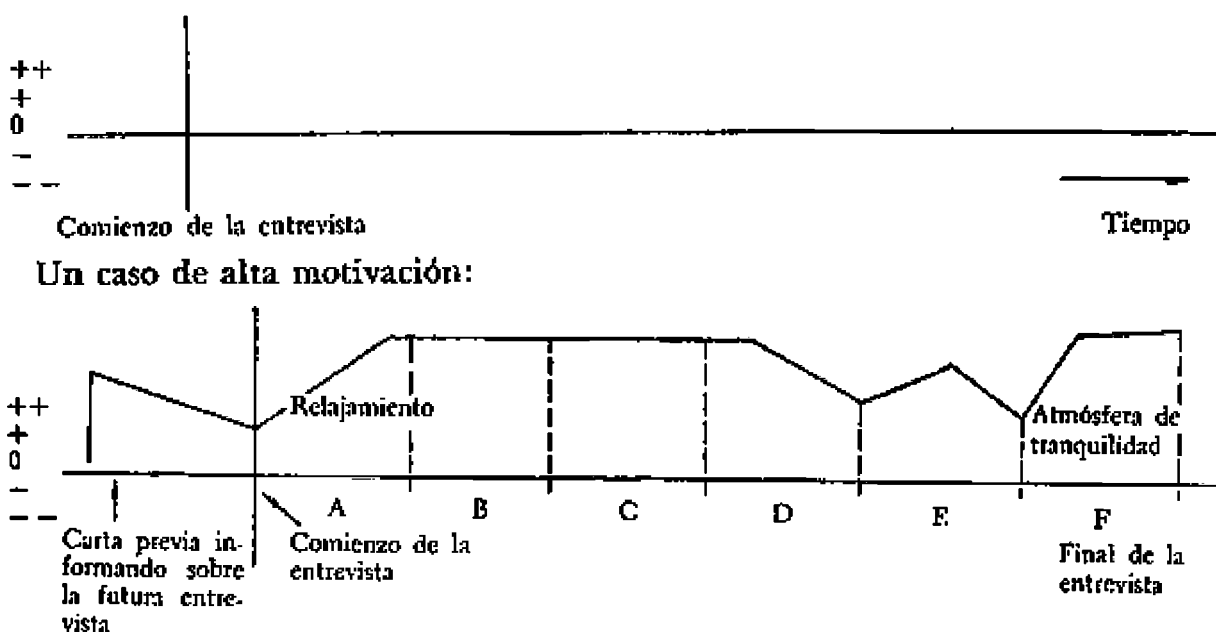
EL OBJETIVO de este capítulo es ofrecer algunos ejemplos de cuestionarios, presentándolos en una forma sistematizada, de tal modo que el investigador posea algunas indicaciones acerca de cómo colocar más fácilmente sus indicadores en el cuestionario. Concentramos nuestra atención en los aspectos más técnicos de la construcción de un cuestionario que pueden contribuir a una minimización de los errores en el registro de los datos.

MOTIVACIÓN EN EL ENTREVISTADO Y EL CUESTIONARIO COMO UNA UNIDAD

El respondiente de una entrevista puede tener un alto o bajo grado de motivación para colaborar y contestar las diferentes preguntas. La motivación depende de factores tales como: objetivo del estudio, quién lo patrocina, el tipo de preguntas, disponibilidad del sujeto, duración de la entrevista, etc., así como de la apariencia y desempeño del entrevistador. Además de estas causas, que se refieren a la entrevista como proceso, hay dos factores de importancia en lo concerniente al cuestionario como unidad, que se refieren a la motivación del respondiente. Éstos son: el orden de las preguntas y el tamaño del cuestionario.

El orden de las preguntas

Estudiemos en un diagrama lo que puede ocurrir durante una entrevista con respecto a estos dos factores. Supongamos que representamos el tiempo de la entrevista en el eje horizontal y el grado de motivación a través del eje vertical:



En este caso, el encuentro ya ha sido preparado para la entrevista, a través de una carta enviada algunas semanas antes del comienzo del trabajo en el campo. La idea de la remisión de una carta previa es que una o varias personas soliciten por escrito al respondente su cooperación, resaltando lo necesaria que es y explicándole, además, el significado del estudio y su utilidad. Si el texto ha sido preparado debidamente y la carta de introducción ha sido suscrita por las personas apropiadas (dependiendo de qué trata el estudio y quiénes son los respondentes), la carta motivará al entrevistado respecto de la entrevista.¹

Cuando llega el entrevistador, encuentra la situación preparada y en muchos casos se enfrenta a un individuo que tiene verdadera curiosidad en saber más acerca de la entrevista. En algunos casos, si la carta de introducción está hecha de un modo poco apropiado, puede ocurrir un efecto contrario al deseado, es decir, un rechazo inicial a la situación de entrevista.

Lo primero que el entrevistador debe hacer cuando la entrevista ha comenzado es lograr que el respondente se relaje y confíe en el entrevistador. Para hacer más fácil esto, necesitaremos un tipo especial de preguntas con que iniciar el cuestionario. Áreas de especial interés para el encuestado, como deportes y otras actividades de tiempo libre, algunas preguntas sobre informaciones acerca de comunicación de masas son, por lo general, adecuadas, si el respondente ha oído sobre este u otro tema, si está interesado en algún artículo de cualquier periódico o revista, etcétera.

En B y C en el diagrama, la motivación no cambia y el tipo de preguntas no es ni mucho ni poco motivador. Por ejemplo: algunas preguntas de hecho, como ocupación, ocupación de los padres, escalas, etc. En D y E nos acercamos al final del cuestionario y se hacen presentes las preguntas emotivas. Por ejemplo, en algunos países, preguntas políticas o de ingresos y otras más personales sobre su familia, amistades, etc. Aquí, normalmente será necesario reforzar la confianza ganada y aun puede ser necesario que

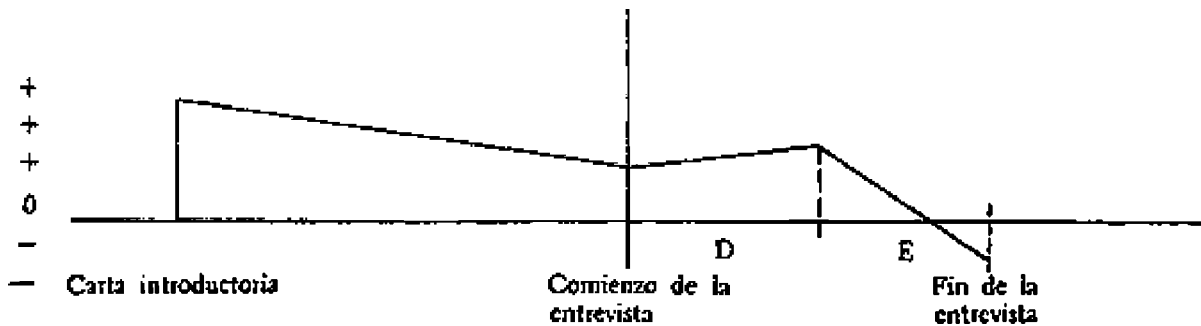
el entrevistador ayude motivando al individuo para una respuesta específica.

En *F* se introduce nuevamente la atmósfera de relajamiento y tranquilidad y es de suma importancia dejar a la persona satisfecha con sus experiencias, aunque debe recordarse que la entrevista siempre se efectúa en forma voluntaria y que el avance de la técnica de entrevista en un país depende de cómo se realicen las entrevistas.

En todos estos pasos, desde *A* hasta *F*, es importante recordar que las preguntas se presentan juntas en campos de intereses similares. Hay una “batería” de preguntas relativas a ocupaciones; otra, a medios de comunicación de masas; una más referente a ingresos y nivel de vida, consumos en la familia, etc. Algunas “baterías” de ítems se presentan como escalas en hojas separadas o en otros métodos especiales que describiremos más adelante. En estos casos es evidente que los ítems representan el mismo campo. En algunos casos el investigador debe repetir un campo de interrogantes antes y después de una batería de preguntas. Esto, naturalmente, con el objeto de encontrar alguna ligazón entre los diferentes campos o, simplemente, para verificar la confiabilidad de las respuestas.

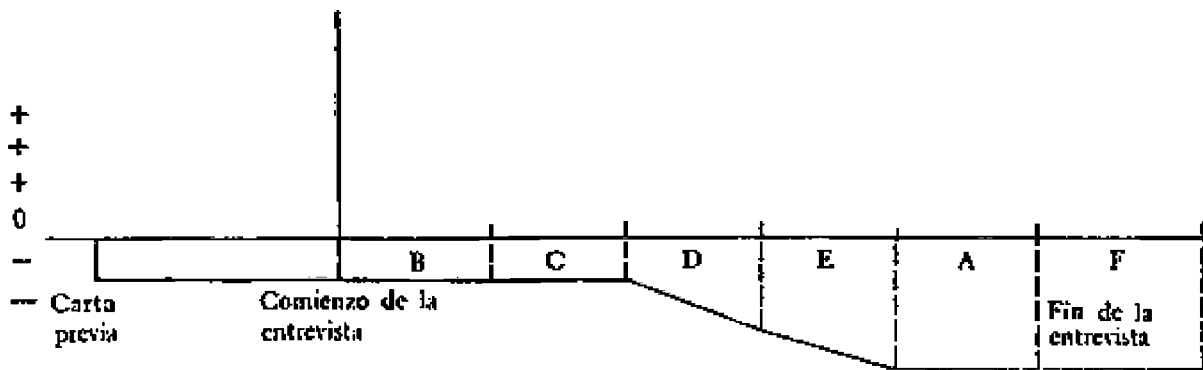
Veamos ahora qué sucede en caso de motivación baja:

Dos casos de motivación baja:



En este caso, el comienzo fue tan favorable como en el ejemplo anterior. Las diferencias están en que el cuestionario parte con el tipo de preguntas *D* y *E*. El respondente se pone suspicaz, se pregunta qué significa todo esto y rehúsa continuar.

Tomemos el caso siguiente:



En este caso el respondente fue muy poco motivado desde la partida, ya que la carta

previa estaba mal hecha y lo puso suspicaz. La entrevista pudo comenzar porque el entrevistador logró motivar al respondiente acerca de su importancia. Sin embargo, el orden de las preguntas no sirvió para arreglar la situación, de tal modo que el *R* se confundió y el resto de la entrevista estuvo a punto de fracasar, aunque llegara a su final natural. Las respuestas dadas a las preguntas fueron cortas, incompletas, sin ningún interés por parte de *R*. En muchos casos, el entrevistador no supo qué contestar y deseaba seguir rápidamente con las siguientes preguntas. Al final de la entrevista, el respondiente estaba seguro de una sola cosa: era la última vez que intervenía en una entrevista.

Aunque la carta previa y la personalidad del *R*, en este caso, fueron factores predominantes para determinar el éxito de la entrevista, un buen enfoque de las preguntas puede, finalmente, dejar en un nivel neutral la entrevista.

Un último comentario referente a una persona “bien motivada” para cooperar en una entrevista. Esta clase de motivación la hemos marcado con ++; indica el máximo de información que el *R* puede dar y no debe confundirse con un estado de sobreexaltación de parte del *R*, que puede ser tan pernicioso como una baja motivación. En este caso, el *R* parte conversando sobre cualquier cosa, sólo por agradar al entrevistador, exagerando hechos e, incluso, posiblemente mintiendo. El procedimiento adecuado que debe seguir el entrevistador en este caso es aquietar los sentimientos, ser más frío, y tratar de concentrarse en actos, dejando al *R* disertando algunos minutos y en seguida concentrarse en el exacto contenido de las preguntas.

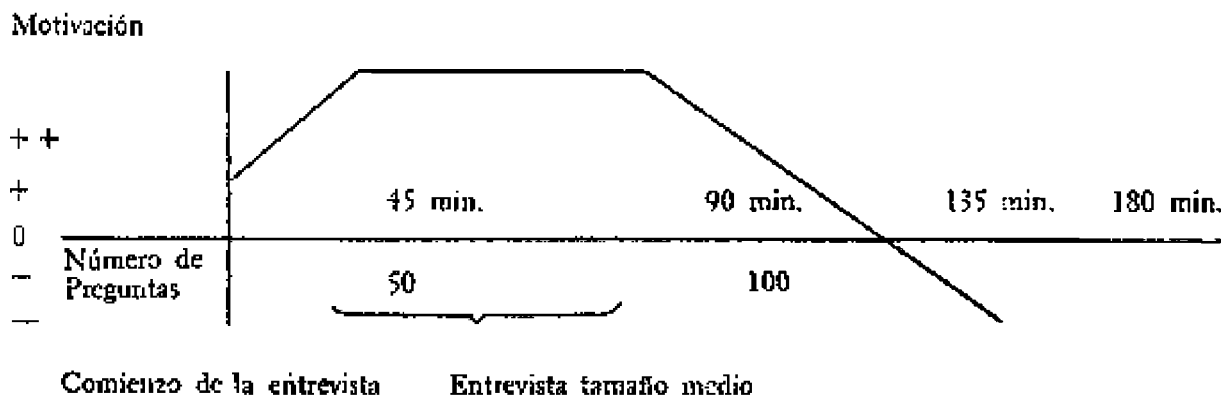
El tamaño del cuestionario

Un cuestionario demasiado corto lleva a una pérdida de información y no da tiempo al *R* para compenetrarse del problema. Existen algunos casos (por ejemplo, entrevistas de mercado e investigaciones de densidades de tráfico) en los que la información que se necesita es limitada y las entrevistas, por consiguiente, deben ser cortas para ser económicamente accesibles.

Un cuestionario demasiado largo, por otro lado, no es aconsejable, ya que podría dificultar mantener en un nivel adecuado el interés del entrevistado. Incluso si se ha informado que la motivación ha sido bastante alta en la mayoría de los casos durante todo el transcurso de la entrevista y que sólo se perdió un porcentaje muy pequeño, no es seguro que el respondiente se preste en otra oportunidad a aceptar una entrevista de tal magnitud. Aunque la entrevista no haya sido traumática para el *R*, le dejó el recuerdo de la pérdida de medio día de trabajo. Por lo tanto el *R* puede negarse a una futura entrevista, aunque en principio dicha persona estuviera interesada, debido al recuerdo del tiempo perdido en la anterior. Hay gastos proporcionalmente excesivos y bastantes dificultades en su manejo en comparación con las ventajas que ofrece sobre un estudio de tamaño mediano. Incluso, el procedimiento de datos y el análisis pueden ocupar tanto tiempo como algunos estudios de tamaño mediano, que aportarían bastante más información a un investigador que un solo estudio de gran tamaño. En algunos casos la

justificación de un estudio extenso se basa en el tipo de análisis a efectuar, por ejemplo, un análisis factorial. Esto es cierto sólo en parte, ya que nuestra búsqueda va dirigida más a la calidad que a la cantidad de las relaciones. Como ejemplo podemos presentar con más utilidad unos 5-6 ítems en forma de escala "Guttman". En ella la escalabilidad ha sido controlada en parte del universo a investigar durante el *pretest*, en vez de presentar un grupo de ítems arbitrarios en un cuestionario final. En el último caso puede suceder que, de unos 30 ítems, sólo 5 o 6 aparezcan realmente bastante correlacionados. Lo que queremos destacar aquí es la necesidad de evitar extender la cantidad de preguntas más allá de los propósitos de la investigación. La multiplicación de preguntas que no se tiene previsto analizar produce pérdidas de tiempo, a veces la extensión del cuestionario invita a los respondentes a rechazar la entrevista, significa costos más altos, y a la larga la mayoría de las veces terminan sin ser utilizadas, ya que después del primer análisis de los datos, los cuestionarios se depositan en cualquier lugar y se olvidan.

Por supuesto, es una labor bastante difícil precisar la amplitud adecuada de un cuestionario de tamaño medio. Esto depende del tipo de preguntas, las que, a su vez, dependen del campo que se está investigando, etc. Trataremos, sin embargo, de dar una primera indicación al lector que está introduciéndose en este tipo de investigación:



Una entrevista con un término medio de 75 minutos toma, en los casos de mayor rapidez, unos 40 minutos, y en los casos excepcionalmente lentos, unos 150 minutos. En algunos cuestionarios hay una parte de las preguntas que se aplican solamente en circunstancias especiales, y otras que se aplican a todos los sujetos. Por ejemplo, las preguntas de la 1 a la 20 pueden estar diseñadas para ser respondidas por todos los *R*; las preguntas de la 21 a la 50 se hacen solamente a los entrevistados que tienen entre 20 y 50 años; las preguntas de la 57 a la 70, a entrevistados mayores de 50 años. Esto naturalmente acorta el tiempo efectivo de la entrevista.

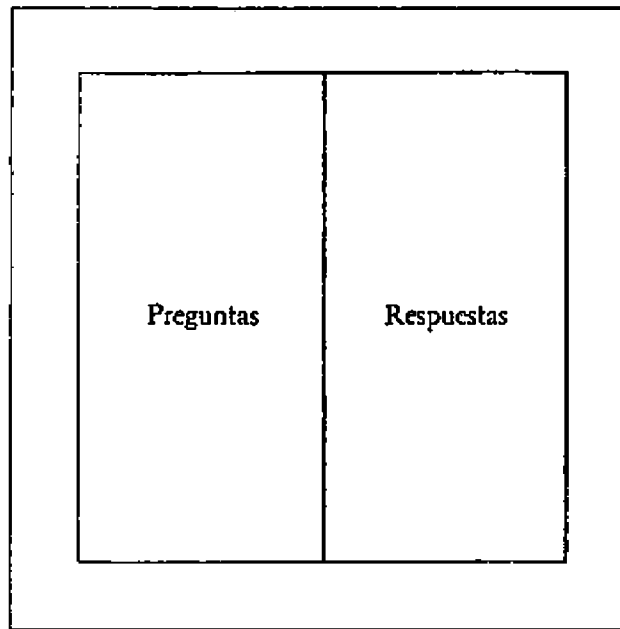
Como indicación general hay que evitar que los entrevistadores se presenten ante los respondentes con un cuestionario que parezca voluminoso (aun en los casos en que el instrumento no se aplique a cada *R* en su totalidad). Conviene en estos casos imprimir el cuestionario con tipos pequeños o separar distintos formularios para diferentes submuestras.

Espacio para las preguntas en el cuestionario

El conjunto total del cuestionario comprende: *a)* un subconjunto mayor en el que se registran las preguntas y las respuestas, y *b)* un subconjunto menor —que ocupa las primeras y las últimas páginas del cuestionario— que contiene información del Registro General de la Unidad, así como comentarios e informaciones de los entrevistadores.

Un ejemplo de la forma de diseñar un cuestionario

Si hacemos un marco en cada página y lo consideramos en dos mitades (una para las preguntas y otra para las respuestas), obtendremos lo siguiente:



Dejemos ahora un espacio separado para la codificación y coloquemos algunas preguntas en las páginas.

Empezamos usando un código general para tres categorías de “respuestas” que son comunes a todas las preguntas:

- O* = Según las instrucciones, esta pregunta no se aplica.
- Y* = Se hizo la pregunta al *R*, pero no quería contestar o no sabía contestarla (en una manera codificable).
- X* = No se hizo la pregunta al *R*, por no estimarlo conveniente el entrevistador.

Esto significa que hemos dejado con alternativas 1, 2, 3, 4, 5, 6, 7, 8, 9, en cada columna de la tarjeta perforada para nuestras categorías de respuestas.

En el *sector núm. 1*, en el ejemplo, hemos colocado *las preguntas*, incluyendo algunas instrucciones fundamentales limitadas (que aparecen subrayadas). Cuando el

cuestionario es administrado por un entrevistador, las instrucciones escritas en el instrumento deben ser mínimas, ya que las pertinentes a cada una de las preguntas deben ser aprendidas en el proceso de entrenamiento. En el caso de investigaciones extensas conviene escribir un *Manual para entrevistadores*, en el que se especifiquen con el mayor detalle características de las preguntas, instrucciones especiales, alternativas, etc. En el caso de cuestionarios autoadministrados, el investigador no debe tener ningún reparo en abundar en las instrucciones.

En el *sector núm. 2* del ejemplo, encontramos *las respuestas*, dadas en forma “cerrada” o “abierta”. En la pregunta 39 el respondente está limitado en su respuesta a tres alternativas. De acuerdo con el resultado de la pregunta, el entrevistador sigue la flecha que le indica cuál es la próxima pregunta que debe efectuar. Si los casilleros se utilizan para colocar la anotación correspondiente a la respuesta, deben preceder al texto de las alternativas. El espacio dejado para la respuesta de una pregunta abierta debe ser lo suficientemente amplio.

	Sector 1 "Preguntas"	Sector 2 "Respuestas"	Sector 3 "Código"
Campo 1	39. ¿Piensa Ud. que sus ingresos y su estándar de vida han aumentado, disminuido o permanecido igual durante los últimos cinco años?	☐ 1 Aumentado → Pregunta 40 ☒ 2 Disminuido → Pregunta 41 ☐ 3 Permanecido → Pregunta 41 igual	43
Campo 2	40. Si han "aumentado" en la pregunta 39, ¿qué es lo que ha causado este aumento? (¿Cuáles son las causas del aumento?)	0	44
Campo 3	41. Si han "disminuido" en la pregunta 39, ¿cuáles han sido las causas de la disminución?	La inflación	45-46
Campo 4	42. Al grupo más joven (nacidos entre 1925-1945). ¿Piensa Ud. que continuará en esta profesión o tiene Ud. planes de cambiarla por otra?	☒ 1 Otra → Preguntas 44, 45 y 46 ☐ 2 La misma → Pregunta 47	47
Campo 5	43. Al grupo de mayor edad (nacidos entre 1890-1925). ¿Qué planes tiene Ud. para su vejez?	0	48

Veamos lo que sucede en la entrevista que se ha explicado a través de nuestro ejemplo:

Para la pregunta núm. 39 fue escogida la alternativa 2. En consecuencia, pasamos a la pregunta 41 y dejamos la 40, que no se aplica, marcada con un 0.

Las anotaciones de *O*, *Y* y *X* pueden colocarse en la línea divisoria entre sector 1 y sector 2. Todo lo que aparezca en el sector 3 se hace después de que la entrevista ha terminado. Si se sigue esta rutina, es muy fácil para el entrevistador y para el investigador que no se pierda ninguna pregunta. Se limita a seguir la línea central y encontrará, sin duda, alguna indicación en cada campo. Puede considerarse como respuesta a una pregunta abierta, tanto una indicación en un casillero o en un *O*, *X*, *Y* o

cualquier indicación escrita. De esta forma, en unos pocos minutos puede verificarse un cuestionario entero y podemos estar seguros de que no nos equivocaremos en este sencillo modo de controlar.

	Sector 1 "Preguntas"	Sector 2 "Respuestas"	Sector 3 "Código"
Campo 1	44. Si "otra" en la pregunta 42, ¿en cuál profesión ha estado pensando Ud.?	Venta de seguros	49.50
Campo 2	45. Si "otra" en la pregunta 42, ¿por qué está Ud. haciendo planes para esa profesión en particular?	Y	51
Campo 3	46. Si "otra" en la pregunta 42, ¿cree Ud. que ascenderá en esa profesión?	Y	52
Campo 4	47. Si "la misma" en la pregunta 42, ¿cree Ud. que progresará más en su actual profesión?	O	53
Campo 5	48. ¿Cuántos aparatos de radio tiene en su hogar?	☐ 1 Ninguno ☒ 2 Uno ☐ 3 Dos ☐ 4 Tres o más	54

Después de la pregunta 41 el entrevistador, automáticamente, pasa a la 42. La respuesta 1 significa que se salta a la 44, pero antes colocará un 0 en el siguiente campo.

En 45 o 46 el R no supo qué contestar y el entrevistador colocó una Y en el campo. Si por alguna razón existe la necesidad de distinguir "Y" en aquellos que no saben contestar y los que rehúsan hacerlo, por ejemplo, en preguntas políticas, el entrevistador

puede perfectamente anotar “Y nq = no quiere” o “Y ns = no sabe” o cualquier otra clave, y esto puede ser perforado separadamente en las tarjetas si encontramos que las diferencias son de interés.

Presentamos a continuación algunas secuencias codificadas de otros R en las preguntas 39-48.

<i>Pre- gunta</i>	<i>Alter- nativa escogida</i>	<i>Pre- gunta</i>	<i>Alter- nativa escogida</i>	<i>Pre- gunta</i>	<i>Alter- nativa escogida</i>	<i>Pre- gunta</i>	<i>Alter- nativa escogida</i>
39	1	39	3	39	Y	39	3
40	Y	40	0	40	0	40	0
41	0	41	0	41	0	41	0
42	0	42	2	42	0	42	0
43	2 (moverse a la ciudad)	43	0	43	0	43	Y
44	0	44	0	44	0	44	0
45	0	45	0	45	0	45	0
46	0	46	0	46	0	46	0
47	0	47	Y	47	0	47	0
48	3	48	1	48	3	48	1

Ejemplo de X en el Código: Si el R es ciego y la pregunta que sigue se refiere a si ve televisión.

El sector núm. 3 en el ejemplo deberá llenarse tiempo después de que la entrevista ha sido efectuada. Este sector posee en cada campo el número de la columna correspondiente marcada en la parte superior izquierda del espacio. La parte de la hoja que contiene el sector 3 puede ser desprendida del cuestionario y entregada directamente al perforador, quien transfiere la información a la tarjeta perforada directamente. Esto normalmente tomará mucho menos tiempo que si la información es transferida primero a tarjetas especiales y desde ellas tomada y enviada al perforador.

ESPACIO PARA EL EMPADRONAMIENTO GENERAL EN EL CUESTIONARIO

El espacio para el empadronamiento general en el cuestionario se hace en la primera página y en la última o en algunas de las últimas páginas. En este espacio el entrevistador y el investigador registran los datos de importancia para la entrevista. Estos datos pueden ser conocidos de antemano, ser obtenidos durante la entrevista o agregados después en la comprobación y proceso de codificación.

La primera página. Veamos un ejemplo de la primera página de un estudio *survey*: (ver grabado p. 81)

Campo núm. 1. En este espacio aparece toda la información necesaria para la localización del respondente en la entrevista, o para una pre-entrevista en el caso de un

“panel”. Cada *R* tiene un número que se le va asignando a medida que se va construyendo la muestra. Este número se repite en el Campo núm. 1 y en el Campo núm. 2. La razón es la siguiente: el investigador puede desprender el Campo núm. 1 de cada cédula de entrevista cuando ésta haya sido realizada. Las personas que van a manejar los datos *a posteriori* (codificadores, por ejemplo) no están ahora en condiciones de identificar al respondente salvo por un número de código. De esta manera el investigador puede tener la seguridad de un alto grado de confidencialidad, sin perder la posibilidad de ubicar a sus individuos en cualquier otra ocasión (para ser entrevistados nuevamente, por ejemplo).

Mientras se trabaje con muestras aleatorias, no estamos interesados en los individuos *per se*, de ahí que se garantice la confidencialidad de los respondentes. En las investigaciones en las que se contempla un “*follow-up*” o en situaciones especiales en las que se necesita individualizar a los respondentes, será necesario obtener algunos datos personales tales como nombre y dirección.

La cuestión de la confidencialidad tendrá que ser estudiada separadamente, ya que es de la mayor importancia para los estudios de este tipo, con el objeto de garantizar un correcto manejo de los datos y de los resultados.

Campo núm. 2. Aquí el entrevistador registra qué ha sucedido con respecto a la entrevista. El entrevistador verifica por sí mismo el cuestionario antes de entregarlo al investigador y agrega sus comentarios. Es muy importante, en investigaciones de gran escala, asegurar que esta operación se realice seria y sistemáticamente.

Campo núm. 3. Contiene el espacio para anotar los asistentes de la investigación que verificaron las entrevistas y el procedimiento de codificación.

Campo núm. 4. Debe anotarse aquí todo el código “general” usado, en forma visible, para que pueda ser consultado fácilmente por el entrevistador.

Las últimas páginas del cuestionario

En estas páginas el investigador debe registrar todas las preguntas que el entrevistador debe contestar de acuerdo con las circunstancias de la entrevista, tales como: la manera en que el entrevistado recibió la entrevista, grado de motivación de la entrevista y grado en que el *R* se encontró en la entrevista y si fue o no perturbada por factores del ambiente en que se desarrolló, etcétera.

CUESTIONARIO

Campo 1	Encuesta nacional de 197... (título de la investigación) Institución Sociológica de	
	Detalles del nacimiento, dirección, nacionalidad, etc. Detalles de la muestra, núms. de registros, etc. Apellidos del R Nombres Entrevistado núm.	
Campo 2	Nombre del entrevistador Núm. id. Entrevistado núm.	
	Tiempo de la entrevista Desde las a las horas. Fecha: 197... Desde las a las horas. Fecha: 197...	
	La entrevista completa Tiempo total de la entrevista: horas minutos	
	Número de llamadas (incluyendo la ocasión en que fue efectuada la entrevista) Incompleta, faltando lo siguiente:	
Campo 3	Entrevista controlada por: Comentarios:	
	Esta entrevista fue hecha el 1, 2, 3, 4, 5, días o después.	
Campo 4	Codificada por: Codificada hasta: Codificación controlada: Comentarios:	
	Código: O = según las instrucciones, esta pregunta no se aplica. Y = se hizo la pregunta al R, pero no quería contestar o no sabía contestarla (en una manera codificable). X = no se hizo la pregunta al R, por no estimarlo conveniente el entrevistador.	
	Institución Sociológica de 197...	

Confidencial

Presentamos un ejemplo de una de estas páginas:

(Estas preguntas son las últimas del entrevistador al entrevistado.)

77. ¿Le agradó contestar las preguntas de este cuestionario?	
78. ¿Se le ha ocurrido ahora algo especial con respecto a la entrevista y que no tuvo oportunidad de mencionar antes?	
79. ¿Había escuchado algo acerca de esta entrevista antes que llegara nuestra carta?	
80. ¿Ha sido Ud. entrevistado en oportunidades anteriores?	☐ Sí. → Pregunta 81. ☐ No. → Fin.
81. ¿Acerca de qué fue entrevistado anteriormente?	

Puede ser recomendable que el investigador agregue algunas páginas en blanco al final del cuestionario, en las cuales el entrevistador presente su “informe” acerca de la entrevista. Esto es: cómo fue recibido y en qué forma continuó la entrevista. Este informe es un documento valioso para el investigador, por las siguientes razones:

- Es una excelente manera de conocer las actitudes del *R* con respecto a la investigación.
- El investigador puede juzgar la calidad del entrevistador (*E*), y controlar casos de falsificación total o parcial de la entrevista.
- Si algunas partes del cuestionario faltan, la razón de ello puede ser encontrada, seguramente, en el informe.
- El informe puede, incluso, servir muchas veces como una ayuda en la categorización de alternativas de respuestas de algunas preguntas abiertas.
- Hacer de cada entrevista un “caso más científico”, acentuando en el *E* la necesidad de un buen entrenamiento teórico y práctico en materia de entrevistas.

Las siguientes preguntas deben ser contestadas por el entrevistador <i>inmediatamente</i> después que la entrevista ha sido llevada a cabo (pero después que el entrevistador ha dejado al R).	
1. ¿El entrevistador y el entrevistado estuvieron solos durante la entrevista?	☐ Sí, estuvieron todo el tiempo en una pieza aislada. ☐ No, otra persona estuvo presente todo el tiempo, pero no interrumpió. ☐ No, otra persona estuvo presente parte del tiempo, pero no interrumpió. ☐ No, otra persona estuvo presente e interrumpió algunas veces.
2. Si el entrevistador y el entrevistado no estuvieron completamente solos. ¿Cree Ud. que las respuestas fueron influidas por la presencia de otra persona?	☐ Sí, hubo influencia sobre varias partes de la entrevista. Partes ☐ Sí, sobre las preguntas ☐ No hubo influencia.
3. Si no fue posible continuar la entrevista. ¿Por qué razón no se pudo continuar la entrevista?	
4. ¿Cuántas visitas realizó al domicilio del entrevistado antes de poder comenzar la entrevista?	La entrevista se realizó en la visita.
5. La actitud del entrevistado ante la entrevista fue:	☐ Muy interesado. ☐ Interesado. ☐ Algo interesado. ☐ Poco interesado. ☐ No interesado.
6. El contacto con el entrevistado fue:	☐ Muy bueno. ☐ Bueno. ☐ Ni bueno ni malo. ☐ Algo malo. ☐ Malo.

— Es un instrumento descriptivo valioso con respecto a algunas actitudes generales. (Por ejemplo, en el caso en que el entrevistador haya anotado aquí lo que *R* tiende a conversar cuando no se le está preguntando.)

DIFERENTES TIPOS DE PREGUNTAS

En esta parte, trataremos de mostrar algunas preguntas estándar y cómo pueden ser introducidas fácilmente en un cuestionario.

La pregunta cerrada

La pregunta cerrada “simple”

Tal como definimos, cualquier pregunta para la cual las posibilidades de respuestas están limitadas a dos o tres alternativas es una pregunta cerrada simple. Un ejemplo de pregunta “cerrada”:

49. ¿Trabaja Ud. actualmente?	Sí. _____ ➔ 50 No. _____ ➔ 52
-------------------------------	----------------------------------

Ventajas. Muy fácil para registrar, interpretar, codificar y analizar, no necesita entrevistadores altamente entrenados.

Desventajas. Su forma impide una clasificación más fina. Las respuestas que se encuentran, justamente, en el borde de dos alternativas tienden a ser forzadas en alguna de las dos categorías. Este tipo de preguntas tiene que ser seguido por aclaraciones en las instrucciones dadas al entrevistador. Por ejemplo: ¿qué se considera como trabajo? Si una persona jubilada trabaja diez horas semanales, pero considera esto como una especie de ocupación de tiempo libre, ¿cómo lo clasificaría en las categorías dadas en la pregunta 49? A menudo al entrevistador se le suscita el siguiente problema: “Yo no sé. He trabajado durante dos semanas, pero esto no es en forma regular”. En este caso, el R tiene que escoger entre si considera que coincide más con el tipo de “trabaja actualmente” o el de “actualmente no trabaja”.

Una generalización como ésta es, muchas veces, muy difícil de efectuar, pero el investigador debe tener claro este tipo de problemas con el fin de homogeneizar las instrucciones a sus entrevistadores, evitando en todos los casos que sean ellos los que tomen este tipo de decisiones.

La pregunta “cerrada” con múltiples respuestas

La pregunta con múltiples respuestas es vulgarmente llamada pregunta de “cafetería” por su semejanza con la situación que se produce en algunos lugares donde el cliente tiene que escoger entre varias alternativas de platillos o bebidas. Un ejemplo de una pregunta tipo “cafetería”:

53. ¿A qué hora del día parte Ud. de su hogar para ir a su trabajo (escuela)? Indique la hora con una precisión de hasta 5 minutos (por ejemplo: 0.55).	☐ 1. Los lunes parto de mi casa a las Los martes, a las Los miércoles, a las Los jueves, a las Los viernes, a las Los sábados, a las Los domingos, a las ☐ 2. Mi horario de salida no sigue los días de la semana (pase a la pregunta núm. 55). ☐ 3. No parto para mi trabajo a una hora especial. El entrevistado es ☐ 4. Realizo mi trabajo en el hogar. El entrevistado es ☐ 5. No tengo ninguna profesión y no voy a la escuela. El entrevistado es ☐ 6. Esposa que es ama de casa. ☐ 7. Otro:
--	---

El *R* tiene que escoger una o más alternativas en una pregunta de “cafetería”. En el último caso, hablamos de “múltiples alternativas” o de “código múltiple”. Esto complica de inmediato el análisis de la situación, requiriendo entrevistadores bien entrenados.

Ventajas. La pregunta de “cafetería” tiene más posibilidades de respuestas que la pregunta cerrada con dos alternativas y por eso da la oportunidad de ampliar la información. Es fácil de codificar y analizar en su simple forma.

Desventajas. En muchas ocasiones es difícil obtener la categoría de respuesta adecuada. En el *pretest*, las preguntas de “cafetería” tienen multitud de veces la forma de preguntas “abiertas” que se analizarán y clasificarán en las alternativas buscadas. Muchas respuestas pueden complicar el análisis de la pregunta.

También es posible desechar algunas alternativas, si la respuesta no cae dentro de las categorías dadas. En este caso, la pregunta podría tener un “final abierto”, esto es, una alternativa “otro” (en el código anterior: núm. 7) y algún espacio para escribir una explicación. Esto acrecienta la calidad de las respuestas pero las hace, a la vez, un poco más complicadas para ser analizadas. En el *pretest* esta solución es muy común, cuando el investigador no está totalmente seguro de las posibles alternativas.

Técnicas especiales: hojas sueltas

Cuando la pregunta de “cafetería” tiene muchas alternativas, lo que provoca dificultad para su lectura y puede ser olvidada por el *R*, el entrevistador puede mostrar una página que incluye dichas alternativas. En especial si a cada ítem le corresponde una serie de categorías que deben ser elegidas, como en el caso de una escala. Las páginas que se muestran al *R* no deben tener ninguna anotación correspondiente al código o que pueda

confundir al *R* en alguna forma. Por ejemplo, que pueda pensar que alguna categoría es mejor que otra. Indistintamente, *puede darse una página a cada uno de los entrevistados*, o bien *la misma puede mostrarse a todos*.

En el primer caso, el *R* puede marcar con un lápiz las alternativas que escoja. El entrevistador lee las alternativas al *R* hasta estar totalmente seguro de que éste ha comprendido perfectamente el procedimiento. Sin embargo, las anotaciones se hacen en la página correspondiente en el cuestionario por el entrevistador. Ésta es la copia exacta de la página que se muestra al entrevistado, pero tiene marcas de código en ella. Las anotaciones que valen son las del cuestionario; la hoja suelta sirve como una ayuda solamente. Aquí hay un ejemplo de un “perfil de intereses”, hecho en forma de escala “Lickert”. La hoja en el cuestionario:

77. Ahora bien, aquí hay otras preguntas sobre sus intereses. Por favor, señale entre las actividades que yo leere cuánto le gusta cada una de ellas.					
	No me gusta	Me gusta poco	Me gusta	Me gusta mucho	Me gusta muchísimo
1. Tomar fotografías	1	2	3	4	5
2. Ir a la iglesia o a otras reuniones religiosas	1	2	3	4	5
3. Ir a ver competencias deportivas	1	2	3	4	5
4. Practicar deportes o hacer gimnasia	1	2	3	4	5
5. Jugar al naípe	1	2	3	4	5
6. Conversar de política	1	2	3	4	5
7. Tocar algún instrumento o cantar	1	2	3	4	5
8. Estar en familia	1	2	3	4	5
9. Reunirse con los amigos o amigas	1	2	3	4	5
10. Bailar	1	2	3	4	5
11. Ir al cine	1	2	3	4	5
12. Leer libros	1	2	3	4	5
13. Tomar parte en reuniones políticas	1	2	3	4	5
14. Trabajar en el jardín o en la casa	1	2	3	4	5
15. Ver televisión	1	2	3	4	5
16. Ir a exposiciones de arte o a museos	1	2	3	4	5
17. Leer revistas	1	2	3	4	5
18. Pescar, cazar	1	2	3	4	5
19. Tomar parte en asociaciones o clubes	1	2	3	4	5
20. Escuchar la radio	1	2	3	4	5

En el segundo caso (la página suelta se mostrará a todos) se recomienda plastificarla, tratando en lo posible que sea pequeña y fácil de manejar. A continuación damos un ejemplo en el cual *R* tiene que poner en orden siete alternativas:

72. ¿En qué sectores cree Ud. que el gobierno debe intervenir más? ¿Podría Ud. ordenar las siguientes tarjetas de sectores colocando en primer lugar aquel sector en que el gobierno debería intervenir más? <i>Ordenar tarjetas sueltas. Poner número de rango.</i>	☐ Educación ☐ Industria ☐ Minería ☐ Prensa ☐ Religión ☐ Salud ☐ Compañías extranjeras
--	--

Se le presentan siete tarjetas pequeñas con alternativas y se le ruega ordenarlas desde la más alta a la más baja.

Religión	Industria	Minería	
Salud	Educación	Prensa	Compañías extranjeras

Esta técnica especial acrecienta la validez y confiabilidad de las preguntas que se efectúan y además constituye una interesante ruptura en la rutina de *R*, que lo acerca al problema resolviendo tareas.

La pregunta abierta

Un ejemplo de una pregunta abierta:

56. ¿Piensa Ud. que hay algo en la sociedad con lo cual Ud. no está de acuerdo, algo que Ud. desea cambiar? <i>(Follow-up. Preguntas "adicionales" o "clarificadoras")</i>	
--	--

En la pregunta "abierta" el número de alternativas de respuestas posibles es prácticamente infinito. En algunas ocasiones *R* tiene mucho que decir y continuará con la

respuesta alrededor de unos 10 minutos; en otros casos, tiene dificultad para decir algo. La habilidad del entrevistador es importante en las preguntas “abiertas”. Casi siempre el entrevistador hace un *follow-up*, es decir, pone preguntas adicionales o preguntas que pueden clarificar más la respuesta. Es muy difícil obtener del *R* que diga más sobre alguna materia sin interiorizarse más en la conversación. Y si el entrevistador tiene que efectuar las preguntas del *follow-up*, tiene que asegurarse primero de maniobrar la técnica de la entrevista, de tal modo que no haga preguntas directivas o preguntas ambivalentes.

Ventajas. Buena comprensión de los motivos y grado de compromiso del *R*. Mejor contacto entre el *R* y el entrevistador. Incremento de la motivación para la entrevista.

Desventajas. En algunas ocasiones, dificultad de parte del *R* para contestar cuando no tiene la respuesta lista. Compare el caso de la pregunta de “cafetería” donde obtiene la ayuda con ciertas clases de respuestas de las preguntas. Dificultad para las anotaciones e interpretaciones. Dificultades para clasificar las categorías para el código. Necesidad de entrevistadores bien entrenados, lo que aumenta el costo del estudio. En algunas ocasiones, los prejuicios del entrevistador pueden influir fácilmente en las respuestas.

La pregunta “abierta” es usada de preferencia en el *pretest* cuando el investigador no conoce con certeza las posibles categorías de respuestas. Pero las preguntas “abiertas” son también usadas en el estudio final, especialmente como las últimas que se colocan en una batería de preguntas para obtener, así, un análisis más profundo.

Las diferentes respuestas de las preguntas “abiertas” deben ser codificadas en un número de categorías, con el objeto de poder analizarlas cuantitativamente. Esto, algunas veces, puede constituir una tarea difícil. La confiabilidad en la manera que codifiquen los codificadores tiene también que ser verificada, en el sentido en que no exista ambivalencia o falta de claridad dentro de las categorías. Esta confiabilidad debe tener un puntaje de correlación mínima de .90 en cualquiera de dos codificaciones realizadas que se tomen. A continuación, presentamos un ejemplo con una pregunta abierta con su código lateral respectivo:

57. ¿Cuáles son, a su juicio, las principales éticas que un profesor contemporáneo debe tratar de desarrollar en sus alumnos? <i>Indique la cualidad más importante.</i>	
--	--

Un ejemplo de codificación de “preguntas abiertas” es el método que se llama “codificación en el campo”. La ventaja aquí reside en que el entrevistador puede estimar directamente los sentimientos acerca de la pregunta y codificarlos en categorías establecidas. Naturalmente, hay que tener mucho cuidado porque es fácil introducir

sesgos por parte del entrevistador. Éste puede equivocarse, por ejemplo, al juzgar al *R* por algunas razones ajenas a la contestación de la pregunta, como su manera de vestir, de hablar, etcétera.

<i>Código Pregun- ta núm.</i>	<i>Colum- na núm.</i>	<i>Alternativas de respuesta</i>
57	62	1. <i>Superación y alegría: espontaneidad, felicidad, curiosidad por saber, creatividad, honestidad, confianza en sí mismo, autonomía</i> 2. <i>Uso racional de la energía: celo, aptitud, tenacidad</i> 3. <i>Habilidad activa para adaptarse a situaciones sociales: sociabilidad, cooperatividad, camaradería</i> 4. <i>Habilidad pasiva para adaptarse a situaciones sociales: auto-control, autodisciplina, comportamiento perfecto, pulcritud, limpieza, tolerancia</i> 5. <i>Habilidad para ajustarse a personalidades autoritarias: respeto por la autoridad, profesores y mayores, subordinación de los sentimientos al deber</i> 6. <i>Ajuste autosustentado a situaciones sociales: deseo de competencia, ambición</i> 7. <i>Amistad, amor, tolerancia, buen corazón</i> 8. <i>Lealtad a los valores existentes: honestidad, sentido de lo bueno y lo malo</i> 9. <i>Importancia personal, valentía, honestidad</i>

Damos a continuación un ejemplo de una pregunta abierta, en la cual el entrevistador, además de escribir la respuesta en el espacio vacío, tiene que hacer una codificación en el campo acerca del grado de relación (favorable-desfavorable) entre profesores universitarios y profesores secundarios. Nótese que el entrevistador *no* lee las alternativas presentadas en el sector de las respuestas.

38. ¿Cuán buena es, en su conjunto, la relación existente entre profesores universitarios y profesores secundarios como grupos profesionales?	☐ Muy buena. ☐ Bastante buena. ☐ Buena y mala. ☐ Bastante mala. ☐ Muy mala.
--	---

La utilización de preguntas cerradas y abiertas, en conjunto.

Ejemplo de un “código múltiple”

Como hemos mencionado en el caso de preguntas tipo “cafetería”, la posibilidad de seleccionar más de una alternativa entre muchas puede complicar el análisis de datos. Este tipo de respuestas las hemos llamado “código múltiple”. La construcción de tales preguntas puede involucrar bastantes problemas.

Presentamos en seguida un ejemplo de cómo se resuelve el problema con alternativas múltiples, utilizando diferentes tipos de preguntas. Supongamos que queremos poner unas preguntas acerca de la percepción de medios de comunicación de masas (“*mass media*”). Empezamos con los diarios y preguntamos acerca de cuál periódico es más leído por los *R*.

Ahora bien, es posible que *R* lea: ninguno, uno, dos o más. Deseamos, asimismo, preguntarle la frecuencia con que lee los diferentes diarios en el caso de que lo haga y, en especial, qué parte del periódico prefiere. Todo esto debe estar dividido en su orden lógico en el cuestionario, con el objeto de que no sea demasiado complicado.

En primer término, ¿sobre cuántos periódicos preguntaremos? Evidentemente, sobre uno como mínimo. Dos periódicos sería mejor todavía, si se trata de una persona que le gusta estar bien orientada en cuanto a noticias y que probablemente lee dos en forma detallada. Existirá una diferencia en percepción de la comunidad, naturalmente, entre las personas que leen uno y las que leen dos diarios, con respecto a sus conocimientos políticos, por ejemplo. Mucho más con alguien que lee y se documenta a través de tres diarios. Podrían, incluso, existir diferencias entre quien lee dos diarios y quien lee tres. Luego seguimos con quienes leen cuatro, etc. Evidentemente, mientras mayor sea el número de diarios escogidos, menor será el porcentaje de *R* a los que atañe la categoría y menor la diferencia en percepción de la comunicación, digamos, entre las personas que leen 5 y las que leen 6 periódicos.

De esta manera, a través de nuestro *pretest* hemos encontrado en una pregunta “abierta” que muy pocas personas leen 3 y más diarios, pero que existe una diferencia apreciable en conocimientos políticos, por ejemplo, entre los que no leen ninguno, y los que leen uno o dos diarios en una proporción más o menos clara. Consecuentemente decidimos preguntar solamente acerca de dos como máximo.

Segundo, ¿de qué manera seguiremos preguntando? Empezaremos por preguntar qué diario lee el *R*; en seguida, qué es lo que lee del diario y, finalmente, con qué frecuencia.

Tercero, ¿qué clase de preguntas colocaremos: abiertas, cerradas, etcétera?

El ejemplo desarrollado en las dos páginas siguientes puede ser una solución.

Hemos presentado una solución al problema de obtener información acerca de la exposición y percepción de medios de comunicación masivos, utilizando los tres tipos de preguntas presentados anteriormente: cerradas simples, cerradas con respuestas múltiples y abiertas.

59. ¿Lee Ud. diarios o revistas?	1. ☐ Sí. 2. ☐ No. —————→ 66
60. ¿Cuál es el diario que lee más?	1. ☐ <i>Excelsior.</i> 2. ☐ <i>Últimas Noticias.</i> 3. ☐ <i>Novedades.</i> 4. ☐ <i>El Sol de México.</i> 5. ☐ <i>El Universal.</i> 6. ☐ <i>El Herald.</i> 7. ☐ <i>El Día.</i> 8. ☐ <i>La Prensa.</i> 9. ☐ Otro. ¿Cuál?
61. ¿Ud. lee también otro diario? ¿Cuál? <i>El que R lee en 2º lugar.</i>	0. ☐ No lee otro. 1. ☐ <i>Excelsior.</i> 2. ☐ <i>Últimas Noticias.</i> 3. ☐ <i>Novedades.</i> 4. ☐ <i>El Sol de México.</i> 5. ☐ <i>El Universal.</i> 6. ☐ <i>El Herald.</i> 7. ☐ <i>El Día.</i> 8. ☐ <i>La Prensa.</i> 9. ☐ Otro. ¿Cuál?
62. Ud. ha dicho que lee el, lo hace:	☐ Todos los días (5-7 días de la semana). ☐ Algunas veces por semana (2-4). ☐ Una vez por semana (1). ☐ Algunas veces al mes (1-3). ☐ Casi nunca. ☐ Otro. ¿Cuál?
63. ¿Qué partes en particular le gusta leer en el? <i>Mencionar primer diario. Follow-up. Código múltiple.</i>	

En general, esta manera de presentar las preguntas evita interpretaciones ambivalentes y las hace más fáciles de registrar, al ordenar los diferentes campos de preguntas, con arreglo a la manera lógica de pensar primero en el área más grande (si lee o no diarios), después preguntando por campos más restringidos y así sucesivamente.

Aunque en el ejemplo es claro que habrá algunas dificultades en la categorización de las preguntas abiertas, podremos no obstante obtener un conocimiento más profundo de qué significa para el R leer diarios. Para algunos R, leer el diario puede significar leer solamente ciertos rubros; para otros puede significar leer en partes o leerlo en su

totalidad.

64. Ud., ha dicho que también lee el; lo hace:	Todos los días (5-7 días de la semana). Algunas veces por semana (2-4). Una vez por semana (1). Algunas veces al mes (1-3 veces al mes). Casi nunca. Otro. ¿Cuál?
65. ¿Qué partes en particular le gusta leer en el? <i>Mencionar Segundo diario. Follow-up. Código múltiple.</i>	
66. ¿Ud. lee revistas?	Si. No. — 74

Ejemplo de preguntas encadenadas

A veces, se puede utilizar una serie de preguntas con el objeto de obtener un conocimiento más profundo de la situación, y también pueden ser usadas para forzar al *R* a clarificar su posición. Presentamos a continuación las siguientes preguntas:

70. Si Ud. tuviera que definir su interés en los asuntos políticos, ¿Ud. diría que está muy interesado, algo interesado, o no muy interesado?	1. ☐ Muy interesado. 2. ☐ Algo interesado. 3. ☐ No muy interesado. 71 4. ☐ Otra respuesta.
71. Si contestó: "No muy interesado" a la pregunta 70: ¿Está usted un poco interesado o no le interesa nada?	1. ☐ Un poco interesado. 2. ☐ Nada interesado.

Muchos respondientes han contestado en la categoría 3 sobre la pregunta 70. En la nueva pregunta, su intención con la respuesta "No muy interesado" es dividirla en dos nuevas categorías: "Un poco interesado" y "Nada interesado". De esta manera hemos conseguido un refinamiento de la escala, en un espacio de la continuación obligada.

La misma idea se ha empleado en la pregunta siguiente (72). Primero, el respondiente ha sido interrogado sobre si está o no de acuerdo. Esto es, un terreno alternativo entre la

alternativa 1 y 2 o 3 y 4. Luego, dentro de cada categoría, ha sido interrogado sobre el grado de aprobación o negación.

De esta forma, hemos omitido un 0 en la escala, en el cual muchos entrevistados tienden a concentrarse, ya que es más fácil dar una respuesta que no constituye una afirmación exacta.

La pregunta entonces sigue y trata primero de explicar los motivos de la actitud y finalmente la importancia que tiene para *R*. Con esto, hemos cubierto la parte más esencial referente a la opinión necesitada.

72. En muchos países en los cuales existe la pena de muerte se discute si debe mantenerse o ser abolida. ¿Cuál es su opinión al respecto? ¿Está Ud. completamente convencido de ello, o sólo piensa que así debe ser en general? <i>(Escriba la respuesta palabra por palabra.)</i>	1. No pena de muerte (completamente convencido). 2. No pena de muerte (convencido en forma general). 3. Pena de muerte (convencido en forma general). 4. Pena de muerte (completamente convencido). Otra respuesta:
¿Podría explicar los motivos de su actitud ante esa pregunta algo más en detalle? <i>(Escriba la respuesta palabra por palabra.)</i>	

Para posibilitar la obtención de toda esta información de manera óptima, hemos utilizado diferentes tipos de preguntas en las baterías. La pregunta 72 tiene, por ejemplo, en su primer paso, el carácter de ser una pregunta cerrada simple. La ventaja en este caso es que existen solamente dos tipos de respuestas y la persona no puede escapar a una categoría media más “conveniente”. El segundo paso de la pregunta 72 es una pregunta tipo “cafetería”, si vemos la pregunta en total globalmente. Aquí obtenemos el grado de aceptación o rechazo.

Finalmente, a fin de obtener las explicaciones de los motivos, hemos utilizado una pregunta abierta para asegurarnos una penetración profunda del problema.

Preguntas particulares

En seguida presentamos preguntas aplicadas a diferentes *campos* de información muy utilizados:

Ejemplo de una pregunta sobre ingresos:

Si vamos a estudiar la distribución de ingresos del país o de la población y nos es desconocida, podremos obtener una idea en el *pretest*, que puede llevarnos a las presentes categorías de respuestas (ver ejemplo). La distribución de ingresos está sesgada hacia el lado de los menores ingresos. Esto tiene que ser representado en las alternativas de respuesta. Una pregunta como ésta debería ser seguida de otras en las que se pregunte cuántas personas viven de ese ingreso, si existen otras fuentes extras de ingreso u otros beneficios. En algunas ocasiones, el jefe de la familia (el único que gana dinero) no coincide con el entrevistado, etcétera.

Estas preguntas debieran ser también elaboradas, de modo que no se presten a confusiones, ya que existen trabajos diferentes que llevan la misma denominación.²

104

Por ejemplo, alguien se define a sí mismo como “pintor”. Puede tratarse de un artista, un empleado o empleador de una fábrica de pinturas. A menudo se puede estudiar la movilidad ocupacional a través de una batería de preguntas de este tipo, agregando preguntas sobre la ocupación del padre y de los hijos. Un ejemplo de preguntas sobre ocupación se ve en la p. 109.

Preguntas sobre conocimientos

Finalmente presentaremos algunos ejemplos de cómo medir el conocimiento (en este caso, el conocimiento político en un cuestionario). Ejemplo de preguntas para verificar conocimientos políticos:

76. ¿En cuál campo se han destacado las personas que se indican a continuación?	Política	Artes	Música	Literatura	Ciencias
1. Béla Bartók					
2. Wernher von Braun					
3. William Faulkner					
4. Alexander Fleming					
5. Heinrich Himmler					
6. Carl Nielsen					
7. Sékou Touré					
8. H. C. Urey					

Los entrevistadores leen las alternativas una por una o muestran una hoja que las contiene. Una respuesta equivocada puede ser anotada con un signo – y una respuesta correcta, con el signo +.

En estas preguntas hay que tener un cuidado extremo de cada detalle: en la pregunta 77, una respuesta correcta se anota con el signo +, y una respuesta equivocada con el signo –. ¡Cuidado! La mitad de los números de los ítems puede ser falsa, es decir, el investigador puede introducir algunos ítems sobre personas o sucesos inexistentes, para probar la confiabilidad de las respuestas.

En la pregunta 78 no se presentan alternativas y el entrevistado tiene que conocer bastante bien el área para poder responder. Tiene la desventaja de que desecha los individuos que tienen algunos conocimientos pero no los suficientes para recordar exactamente las respuestas correctas. Una dificultad estriba también en la interpretación de algunas respuestas medio correctas.

77. ¿De cuáles de los sucesos siguientes se ha enterado Ud.? Mostrar hoja suelta. (Ninguno sucedió antes de 1973.) Marque con una X en el cuadro indicado.		
	<i>Me he enterado</i>	<i>No me he enterado</i>
1. Nixon renuncia a la presidencia de los Estados Unidos.	☐	☐
2. La India explota su primera bomba atómica.	☐	☐
3. China Comunista compra trigo a Canadá.	☐	☐
4. Asesinan a Carlos Pratts en Argentina.	☐	☐
5. Escándalos de espionaje en Alemania del Oeste causan la renuncia de Willy Brandt.	☐	☐
6. Juan Perón muere y su esposa lo sucede en la presidencia de Argentina.	☐	☐
7. Cactano derrocado por los militares portugueses.	☐	☐
8. El huracán Fifi destruye la economía de Honduras.	☐	☐
9. Los países productores de petróleo suben los precios del crudo.	☐	☐
10. El director de la CIA declara que USA interviene en el derrocamiento del presidente Allende, en Chile.	☐	☐
11. Valéry Giscard D'Estaing es elegido Presidente de Francia.	☐	☐

78. 1. ¿Podría Ud. decir el nombre del secretario de Educación?	
2. ¿Conoce Ud. el nombre del líder del Partido Acción Nacional?	
3. ¿Conoce Ud. a cuál partido político representa el presidente de los Estados Unidos?	
4. ¿Podría nombrar el partido político al cual los siguientes periódicos representan o con el cual simpatizan? <i>Excelsior, El Universal, El Día.</i>	
5. ¿Los siguientes nombres, le significan algo a Ud.? Carlos Andrés Pérez, Velasco Alvarado, Hugo Vanzer.	
6. ¿Ha visto Ud. alguna vez las siguientes abreviaturas y podría Ud. mencionar el significado que tienen? O.I.T., NAFINSA, C.T.M.	
7. ¿Sabe Ud. cuántos diputados tiene el Congreso?	

Las preguntas 76 a 78 se presentan para ser aplicadas a los estudiantes universitarios. No puede aplicarse a la población en general pues son bastante difíciles.

“LISTA DE CONTROL” PARA CUESTIONARIOS Y ENTREVISTAS

En las siguientes líneas tratamos de presentar tres “listas de control”:³ una para la construcción del cuestionario, una para el cuestionario y las baterías de preguntas y una última que se refiere al cuestionario y a la entrevista.

Una “lista de control” para la construcción del cuestionario

1. Escriba el texto de la pregunta lo más simple posible. La pregunta puede ser muy comprensible para la persona que la escribe. En muchas ocasiones el investigador emplea un lenguaje demasiado complejo en su tarea diaria y en su conversación profesional con el entrevistado. Intente hacer más accesible el lenguaje dejando de lado las palabras abstractas, el “lenguaje científico” y términos como “nivel de educación”, “movilidad”, “percepción”, etcétera.

2. La pregunta no debe contener más de 25 palabras. Mientras más corta, mejor, siempre que su contenido no sufra. La pregunta debe expresarse de tal modo que no necesite ninguna explicación adicional de parte del entrevistador.

3. La pregunta debe contener una sentencia lógica. Si mezclamos dos sentencias lógicas, la pregunta es ambigua. Ejemplo: “Los profesores son justos con los estudiantes que están de acuerdo con sus ideas e injustos con los estudiantes que no lo están”. Un acuerdo con esta sentencia puede implicar acuerdo con dos juicios o acuerdo con uno de ellos que parezca predominante, pero no sabemos con cuál. Es recomendable convertir en dos preguntas ésta que acabamos de mencionar en el ejemplo.

4. Evitar las preguntas ambiguas. Ejemplo: “Las mujeres son capaces de todo”, puede tomarse en el sentido positivo o en el negativo que encierra la frase.

5. Evitar el uso de estereotipos o de palabras cargadas emocionalmente. Ejemplo: “La mujer es sólo una variante inferior del género humano”.

6. Evitar preguntas dirigidas. Esto es, preguntas cuya respuesta esté sugerida al entrevistado. Ejemplo: “a) La razón de la sequía es la falta de agua; b) ¿No cree usted que existe un alto grado de inflación en este país?”

7. El investigador debe decidir si personificará las preguntas o no. Esto depende de la gente con la cual va a trabajar. En la mayoría de los casos es recomendable personificar las preguntas. Cuando el contenido de la pregunta no es socialmente aceptado, la pregunta no debe ser personificada, ya que puede conducir a un rechazo.

8. Decidir en qué sentido, deseado o no, modificará las respuestas el empleo de nombres de personajes. Se ha visto que, en la medida en que el sujeto no conoce la opinión del personaje sobre el problema planteado, responderá según la simpatía o antipatía que éste le cause. Ejemplo: “¿Aprueba o desaprueba usted las ideas de Carlos Marx sobre los obreros?”

9. Para el fenómeno de la deformación conservadora, una tendencia a la desconfianza o temor al cambio:

Para la desconfianza. Enunciar las preguntas de manera tal que los acuerdos o desacuerdos no estén en un mismo nivel respecto de la opinión. Vale decir que a veces estar en favor de algo implique rechazar la sentencia, y otras aceptarla. Ejemplo: “Los capitalistas en conjunto son un grupo de gente decente”.

**Completamente
de acuerdo**

De acuerdo

Indeciso

**En
desacuerdo**

**Completamente
en desacuerdo**

Para el temor al cambio. Deben formularse las preguntas de manera tal que, si se proponen modificaciones de las normas vigentes, sea aparentemente menos violento.⁴ Ejemplos: “¿Es usted partidario de *interpretar* legalmente la Constitución de manera que se impida...?” ¿Es usted partidario de *modificar* la Constitución de manera que se impida...? En este caso, la segunda formulación es más atrayente para un rechazo proveniente del temor de reconocer una tendencia al cambio. La primera es preferible.

10. Evitar una pregunta que parezca poco razonable o inoportuna, empleando una breve justificación introductoria de por qué la incluimos.

11. Evitar preguntas que requieren mucho trabajo de parte del informante, para *evitar la fatiga*, y también, a veces, errores de memoria. Ejemplo: “¿Cuáles han sido sus ingresos mensuales desde 1950 a 1970?”

12. Ligado a lo anterior, si nuestra pregunta tiene una larga lista de alternativas (pregunta de “cafetería”), para el caso de la entrevista (respuesta oral a estímulo auditivo), es preferible acompañar la pregunta con la lista de respuestas alternativas en una tarjeta que el informante pueda tener presente.

13. En general, sin embargo, es preferible no dar largas listas de alternativas, y aún menos si éstas son *difíciles* de recordar durante la entrevista. Si se le dice: “Ante esto, usted prefiere... 1..., 2..., 3..., 4..., 5..., 6..., 7..., etc.”, cuando acabamos de leer la lista, el sujeto ha olvidado nuestras alternativas. *En realidad estamos jugando con datos cargados.*

14. Si nuestra pregunta contempla un ordenamiento de las alternativas ofrecidas, es preferible imprimir cada alternativa en tarjetas separadas que el sujeto pueda manipular y ordenar a su gusto.

15. Para las preguntas abiertas, deberá dejarse amplio espacio con el objeto de anotar las respuestas.

16. En general, *las alternativas a nuestras preguntas* (en el caso de emplear este sistema) *deben ser formuladas en base a un análisis de las respuestas dadas por los sujetos del pretest* cuando se les preguntó sobre el tema, sin ponerles ninguna restricción, de manera abierta. *Extraeremos las alternativas de entre las respuestas espontáneas más frecuentes y realistas.* En suma, *la formulación de las preguntas debe garantizarnos el logro de los objetivos básicos en el empleo de este tipo de estímulos en la investigación:* a) *Validez:* que las preguntas den cuenta del fenómeno que estudiamos; b) *Discriminación:* que nos permita discriminar subgrupos entre los R; que obtengamos respuestas en diferentes sentidos; si todos coinciden, poco uso le daremos al ítem.

Por último, recordaremos nuevamente que, si entre los datos o variables que queremos estudiar anotamos *actitudes* (ideológicas y otras variables de un nivel de organización superior al de las simples opiniones), *debemos recurrir en lo posible a instrumentos escalares;* si anotamos *actitudes, motivos, variables de personalidad,* recurrir entonces a las técnicas autorizadas, teniendo en cuenta que si no son practicables en un *survey* masivo no hay suficiente razón para remplazarlas por *preguntas de*

cuestionario.

Una “lista de control” para el cuestionario final como un todo y para las diversas baterías de preguntas

1. Las preguntas deberán agruparse de modo que formen una unidad (batería).
2. Las preguntas deben estar ordenadas de modo tal que exista una progresión lógica en la entrevista, de manera que: *a)* el informante sea introducido en la entrevista despertando su interés; *b)* pase de las preguntas más simples a las más complejas; *c)* se procure no enfrentar al informante con un pedido prematuro y súbito de datos personales; *d)* aunque no esté directamente relacionado con la formulación de las preguntas, debe anotarse que, cuando se efectúa una pregunta que puede provocar embarazo en el entrevistado, ha de dejársele oportunidad para que se explique, con el objeto de no frustrar la eficacia del resto de la entrevista, y *e)* se conduzca el interrogatorio de un marco de referencia a otro lo más suavemente posible sin efectuar saltos bruscos.
3. Las preguntas introductorias del cuestionario deberán ser atrayentes sin provocar controversias. Se aconseja comenzar con cuestiones aparentemente irrelevantes, inofensivas o neutras, que despierten simultáneamente los intereses del *R* y “lo introduzcan” en la entrevista.
4. Las preguntas más cruciales o estratégicas del instrumento no deberán aparecer al principio de la entrevista —por las razones ya vistas— ni al final, cuando la fatiga puede influir en el ánimo del *R*.
5. Las preguntas que se refieren a aspectos íntimos del respondente deberán dejarse para el final de la entrevista, cuando ya existe presumiblemente mayor confianza entre entrevistador e interrogado.
6. El cuestionario debe finalizar con expresiones de agradecimiento por la colaboración prestada por el entrevistado, lo que facilitará el contacto para eventuales entrevistas futuras, de la misma o de otra investigación.
7. Si posteriormente es necesario para el relevamiento separar rápidamente a los sujetos de la muestra en grupos previamente determinados, se recomienda emplear papeles de distintos colores (o bandas de colores diversos en los extremos superior derecho o izquierdo del documento) para cada grupo.
8. Se recomienda que las preguntas de una batería vayan desde los tópicos más generales a los más específicos.
9. Para los temas principales de la investigación, planear incluir preguntas que ubiquen el mismo contenido en diferentes contextos, con el fin de poder comparar luego las respuestas.
10. Si hemos empleado una batería de ítems para algunos temas, es conveniente incluir alguna(s) pregunta(s) que nos sirvan para verificar la adecuación y consistencia de las preguntas como un todo.
11. Si planificamos emplear el cuestionario posteriormente (por ejemplo, para

estudios *cross-cultural*) debemos redactar las preguntas de manera que no se hagan temporales; es decir, que se puedan aplicar después sin necesidad de introducir cambios, para asegurarnos las posibilidades de comparación. Esto significa evitar el empleo de fechas, nombres de personajes efímeros, modas o normas locales, etcétera.

12. Si nuestro estudio va a compararse con otros realizados por autores distintos, debemos emplear, en lo posible, la misma formulación de las preguntas.

“Lista de control” de asuntos de entrevistas relacionadas con el cuestionario

Preguntas estructuradas. Regla fundamental: La pregunta debe ser leída en el lenguaje *exacto* con el que figura en el cuestionario. Reglas suplementarias:

1. Tratándose de una pregunta de hechos (como edad, estatus social, etc.), el entrevistador tiene libertad para reformular la pregunta. Lo importante es que la información que deseamos obtener se consiga efectivamente.

2. Si el entrevistado no ha comprendido la pregunta o, sencillamente, la malinterpreta, el entrevistador está *autorizado* para *repetirla*. Sin embargo, no está autorizado para darle otra formulación.

3. Las preguntas *follow-up* tienen que ser hechas por el entrevistador. Incluso si el respondente ha contestado en buena forma, el entrevistador tiene que asegurarse de que el entrevistado no tiene nada más que agregar. El trabajo de las respuestas agregadas a la pregunta *follow-up* debe ser anotado por el entrevistador.

4. En general, si una pregunta es algo débil o completamente irrelevante y su repetición (núm. 2) no da ningún resultado, el entrevistador puede tratar de hacer una pequeña reformulación. En este caso, la respuesta original debe ser anotada, y con el fraseo de la nueva formulación (palabra por palabra) obtendremos la nueva respuesta dada por el entrevistado, la cual debe separarse de la primera respuesta.

5. Si existen alternativas fijadas para la respuesta (pregunta de “cafetería”), frecuentemente sólo se anota *una* de las alternativas.

6. En el caso de una pregunta abierta, la respuesta exacta del respondente debe ser dada palabra por palabra. Por ejemplo: *No*; “el R. piensa que... pero”, *sino*: “Pienso que no es correcto que...” Cualquier desviación de las reglas dadas debe agregarse a las instrucciones de los entrevistadores.

Para *preguntas no estructuradas* las reglas deben aplicarse en cada caso por separado.

UN EJEMPLO

A continuación presentamos un cuestionario autoadministrado, utilizado para la evaluación de un programa de alimentación escolar en cuatro provincias chilenas (Santiago, Valparaíso, Aconcagua y O’Higgins), en 1967.⁵

Se aplicaron 32 500 cuestionarios en una muestra aleatoria estratificada de escuelas, en las que se solicitó información de alumnos y familiares. Los cuestionarios fueron distribuidos por los maestros, y el total de la información se recogió en un mes.

El cuestionario que sirve de ejemplo es el de jefes de familia, y el número en la

esquina superior izquierda es idéntico al del cuestionario del niño, que fue completado por sus maestros, y a los cuestionarios de padres y madres de familia. Para asegurarnos de la complementación entre los dos diferentes cuestionarios solicitamos al maestro que escribiera el nombre del alumno en cada cuestionario.

El cuestionario es además precodificado, salvo en una pregunta (la núm. 24), y los números en la última columna corresponden a columnas de tarjetas IBM (las primeras 43 columnas contenían la información correspondiente al cuestionario de niños; y las últimas 8 a características de la escuela).

Los dibujos en la parte correspondiente a las instrucciones resultaron muy útiles en términos de motivación de los respondentes (el porcentaje de retornos fue de 96.5 % del total de la muestra).

CUESTIONARIO PARA JEFES DE FAMILIA

¡LEER ESTO ANTES DE EMPEZAR!

Aquí tenemos algunas preguntas que queremos que usted responda.

Jefe de familia es aquella persona que tiene la mayor responsabilidad económica en el hogar; si usted no es jefe de familia, se ruega dar este cuestionario a esa persona para que lo llene. (Si no está presente en los 7 días que siguen, llénelo usted en su nombre.)

Las preguntas que usted conteste serán usadas para un estudio científico, por lo tanto son confidenciales y anónimas. El presente estudio es realizado por la Facultad Latinoamericana de Ciencias Sociales (FLACSO,) organismo de la UNESCO (Naciones Unidas). Cuenta con la colaboración del Ministerio de Educación, de la Junta Nacional de Auxilio Escolar y Becas y de la Cátedra de Nutrición de la Escuela de Salubridad de la Universidad de Chile.

Para el propósito del estudio es sumamente importante su colaboración y que responda a todas las preguntas en forma verdadera y lo más exacta posible.

Cuando esté listo el cuestionario, devuélvalo lo más pronto que pueda a la escuela.

Si usted encuentra algunas dificultades en las preguntas, le recomendamos acudir al profesor de su niño para que le ayude.

Ahora: ¿Cómo se responden las preguntas?

- 1.—Marque sólo una respuesta en cada pregunta.
- 2.—Marque las respuestas en la columna que dice “respuestas”.
- 3.—¿En qué forma marco yo las respuestas?
 - a) Si a la respuesta corresponde un número, márkelo con un círculo alrededor del número.

Aquí hay un ejemplo de esto: → La pregunta es ¿Dónde vive usted?

En el campo 1
 En una ciudad pequeña 2
 En una gran ciudad 3

Cuando viva en el campo, por ejemplo, marque como se indica arriba.

b) Si se encuentra con una línea que tenga esta forma ———— coloque en ella la respuesta.

Un ejemplo: ¿Cuántas personas viven en la casa? ————

Si viven 5 personas en la casa, marque 5 ; si viven 12 personas, marque 12 .

4.—¿Cuándo yo hago un error?

En caso de que usted se haya equivocado en una respuesta, borre o haga una cruz grande en la respuesta incorrecta y marque luego la respuesta correcta.

Este es un ejemplo: 1
 2
 3

Agradeciéndole desde ya su valiosa colaboración, quedamos a su entera disposición para cualquier consulta.

FLACSO

J. M. Infante 85 — Santiago

CUESTIONARIO PARA LA FAMILIA DEL ALUMNO

Por favor, ponga el nombre del alumno aquí:

Para comenzar quisiéramos algunos datos sobre usted. (Si usted no es el jefe de familia, tiene que contestar las preguntas en el nombre de él.) Respuestas No marque nada aquí

PREGUNTAS

1.—Sexo del jefe de familia. (Redondee con un círculo el número que corresponde.)

Hombre	1	44
Mujer	2	

2.—Edad de usted (jefe de familia), marque así: Por ejemplo, si usted tiene 30 años, marque 30 en la línea de al lado, si usted tiene 56 años, marque 56

	_____	45-46
--	-------	-------

3.—Parentesco que tiene usted (jefe de familia) con el niño.

Es el padre	1	
Es la madre	2	47
Otro pariente	3	
No es pariente	4	

4.—¿Cuánto gana usted (jefe de familia) mensualmente? Marque así: si gana, por ejemplo, E° 240 mensual (o sea, E° 60 semanal), redondee el número 2.

Menos de E° 100 mens.	o	Menos de E° 25 sem. ...	1	
De E° 101 a E° 300 mens. o	De E° 26 a E° 75 sem.	2		
De E° 301 a E° 500 mens. o	De E° 76 a E° 125 sem.	3		
De E° 501 a E° 700 mens. o	De E° 126 a E° 175 sem.	4	48	
De E° 701 a E° 1 000 mens. o	De E° 176 a E° 250 sem.	5		
De E° 1 001 a E° 1 500 mens. o	De E° 251 a E° 375 sem.	6		
De E° 1 501 a E° 2 000 mens. o	De E° 376 a E° 500 sem.	7		
De E° 2 001 a E° 3 000 mens. o	De E° 501 a E° 750 sem.	8		
De E° 3 000 o más mens.	o De E° 751 o más sem. ..	9		

5.—

PREGUNTAS

	Respues- tas	No marque nada aquí
¿Qué estudios ha realizado usted? (jefe de familia), redondee con un círculo el número o letra que corresponde.	Mar- que aquí	
No fue a la escuela	1	
Fue hasta 2º año de escuela primaria o menos	2	
Fue hasta el 4º año de escuela primaria	3	
Fue más de 4º año, pero no llegó a completar 6º año	4	
Completó la escuela primaria	5	49
Estudios secundarios incompletos	6	
Estudios secundarios completos	7	
Otros estudios posteriores al secundario, pero no de carácter universitario	8	
Estudios universitarios incompletos	9	
Estudios universitarios completos	X	
Otros estudios, especificar	Y	

6.—Total de dinero que entra en la familia.

Indique aquí cuánto dinero reciben los miembros de la familia en total. Sume los sueldos y salarios de todos los que trabajen en la casa. Coloque un círculo en el número que corresponde. Por ejemplo, si ganan en total E° 400, o sea, E° 100 semanal, redondee el número 3.

Menos de E° 100 mens.	o	Menos de E° 25 sem.	1	
De E° 101 a E° 300 mens. o	De E° 26 a E° 75 sem.	2		
De E° 301 a E° 500 mens. o	De E° 76 a E° 125 sem.	3		
De E° 501 a E° 700 mens. o	De E° 126 a E° 175 sem.	4	50	
De E° 701 a E° 1 000 mens. o	De E° 176 a E° 250 sem.	5		
De E° 1 001 a E° 1 500 mens. o	De E° 251 a E° 375 sem.	6		
De E° 1 501 a E° 2 000 mens. o	De E° 376 a E° 500 sem.	7		

De E° 2 001 a E° 3 000 mens. o De E° 501 a E° 750 sem.	8
De E° 3 001 a E° 5 000 mens. o De E° 751 a E° 1 250 sem.	9
Más de E° 5 000 mens. o Más de E° 1 251 sem.	X

7.—¿Cuál es su ocupación? (la del jefe de familia).

Por favor, redondee con un círculo el número que corresponde al grupo de su ocupación. Antes de hacerlo, lea detalladamente todos los grupos y, cuando no encuentre escrita la suya, marque el grupo de ocupaciones que más se le parecen.

Estoy en el grupo de:	<i>Trabajadores no especializados, asalariados y por cuenta propia.</i> —Por ejemplo: ayudantes, aprendices, industriales de construcción y minería, garzones, estibadores, porteros, propietarios de pequeña parcela rural que trabajan la tierra por su cuenta y sin empleados, vendedor ambulante, mozo de servicio y limpieza, diariero, lustrabotas, empleada doméstica, basurero, peón agrícola, pololero y similares	1	
Estoy en el grupo de:	<i>Trabajadores especializados, asalariados y por cuenta propia.</i> —Por ejemplo: técnicos no titulados, linotipistas, gasfitero, carpintero, estucador, pintor, mecánico, operador de máquinas, peluquero, chofer, bencinero, pescador, minero, soldado, cabos, marineros y similares	2	
Estoy en el grupo de:	<i>Personal de supervisión.</i> —Por ejemplo: capataces de industria, minería y construcción, suboficiales de las fuerzas armadas, superiores a cabos (sargento, suboficial, suboficial mayor), jefe de cocina y similares	3	
Estoy en el grupo de:	<i>Empleados de administración pública, privada y comercio de grado II y comerciantes minoristas.</i> —Por ejemplo: oficinistas de administración pública y privada, empleados de comercio, vendedor de mostrador, viajante, corredor de seguros, dactilógrafos, telefonistas, comerciantes minoristas con un empleado o ninguno, y similares	4	51
Estoy en el grupo de:	<i>Empleados de administración pública, privada y comercio de grado I y pequeños propietarios de comercio.</i> —Por ejemplo: técnicos titulados, secretarías, enfermeras, fotógrafos, dibujantes, laboristas, profesores primarios, propietarios de		

	pequeña empresa, o comercio minorista con 2 hasta con 5 empleados y similares	5
Estoy en el grupo de:	<i>Jefes menores de administración pública y privada.</i> —Por ejemplo: jefes de grupo o de oficina, profesores secundarios, oficiales de las fuerzas armadas, desde subteniente hasta capitán inclusive, director de escuela primaria y similares	6
Estoy en el grupo de:	<i>Empresarios medianos y jefes intermedios de administración pública y privada.</i> —Por ejemplo: dueños de empresa o comercio, con 6 hasta con 30 dependientes, secretarios de facultades, o rectores de colegio secundario, jefes de sección o departamento, jefes de fuerzas armadas, de mayor a coronel, profesionales universitarios que ejerzan su profesión y similares	7
Estoy en el grupo de:	<i>Grandes empresarios y altos ejecutivos.</i> —Por ejemplo: dueños de grandes empresas industriales, agrícolas y comerciales con más de 30 empleados, directores de escuelas universitarias, decanos, rector, director general, subsecretarios, ministros, gerentes directores, presidentes y vicepresidentes de grandes empresas, generales, almirantes y similares	8
	Escriba ahora, con algún detalle, en qué consiste su ocupación (del jefe de familia), qué es lo que hace	
		
		

V. LA CODIFICACIÓN *

HÉCTOR APEZECHEA

ANTES DE considerar el tema propio de este trabajo, conviene hacer unas breves puntualizaciones para situarlo dentro de los límites que nos hemos propuesto, y señalar algunos de los alcances de la problemática de la codificación que no tienen cabida lógica en el cuerpo del trabajo.

En primer lugar, se trata de estudiar el problema de la codificación en las investigaciones sociológicas. Su utilización en otras ciencias sociales —demografía, psicología, ecología, geografía humana, etc.—, que el autor desconoce, podría asumir otras variaciones y alternativas de empleo. Por lo tanto no se mencionarán en las líneas que siguen.

En segundo lugar —y salvo las menciones que se hacen en el apéndice—, procuraremos hablar de la codificación en aquellas investigaciones que utilizan como instrumento de relevamiento el cuestionario estructurado, y que tienen como objetivo elaborar los datos mecánicamente con tarjetas perforadas tipo IBM, lo cual, por cierto, restringe en buena medida el campo que abarca este trabajo.

En tercer término, cabe señalar que, en el marco del diseño total de una investigación, los problemas que suscita la codificación de algún modo “invaden” la casi totalidad del diseño. Así, es razonable pensar que, aunque la etapa propiamente de codificación se ubica simultáneamente o posteriormente al relevamiento —bien sea al proponer el elenco de “variantes” (variables operacionalizadas) y de indicadores de las mismas, al confeccionar el instrumento de relevamiento o al proyectar el plan de explotación de los datos—, se tengan en cuenta las ventajas y obstáculos que el procedimiento de codificación trae aparejados. Es una cuestión dudosa, a mi entender, que la codificación en cuanto tal imponga limitaciones tanto a la proposición de la o de las hipótesis a probar empíricamente como a las variables o sus respectivas dimensionalizaciones. En tal sentido, sería solamente al nivel de las variantes cuando comenzaría a tener validez el hecho de que la codificación debe ser tenida en cuenta. No parece necesario ejemplificar los asertos contenidos en este párrafo. En el curso de la exposición que sigue el lector, por sí mismo, podrá percibir en qué consisten las limitaciones y ventajas del procedimiento de codificación y de ella en sí misma, así como cuáles deberán ser los elementos que el investigador deberá tomar para que este procedimiento agilice y no entorpezca el desarrollo de la contrastación empírica de sus hipótesis.

La codificación

La codificación es, genéricamente, la primera operación a realizar en el marco de la elaboración de datos, inmediatamente después de efectuado el relevamiento. (Si bien más adelante veremos que puede comenzar simultáneamente con el relevamiento, como es el caso, por ejemplo, del cuestionario de preguntas cerradas.) Conviene comenzar por las definiciones que algunos autores han dado, para tener una idea de qué se entiende por codificación en sociología y extraer las características más relevantes asignadas a esta operación.

Según Maurice Duverger,¹ tal operación consiste en asignar un número a cada categoría de respuestas, el cual determinará el lugar de la correspondiente perforación sobre la ficha que sirve de base a los escrutinios mecanográficos.

Hyman explicita más el sentido de la operación,² al aclararnos que es un proceso de clasificación dentro de un limitado número de categorías de las respuestas de cada persona. Su propósito es resumir los datos originales de manera tal que puedan ser transferidos a las tarjetas perforadas y tabulados mecánicamente.

Por último, nos parece importante hacer referencia a Seltiz³ y otros autores —más adelante veremos por qué—, cuando nos dice que es el procedimiento técnico por el cual los datos originales son transformados en símbolos (generalmente números) de modo que puedan ser tabulados y contados. La transformación no es, sin embargo, automática: envuelve juicio de parte del codificador.

Podemos tener una noción aproximada de la codificación a través de las citas,⁴ que en parte se repiten y en parte se complementan, pero nos parece importante destacar un aspecto de la definición de Seltiz: la codificación implica la transformación de los datos relevados en *símbolos*,⁵ en general numéricos. Esta noción parece darnos la clave fundamental acerca de la naturaleza de la codificación: si bien es cierto que ésta y la categorización de los datos son dos operaciones que en la práctica se realizan simultáneamente —pudiendo confundirse en una sola—, conceptualmente son dos aspectos que deben manejarse por separado.

La codificación, considerada como operación aislada, consiste en la atribución de un determinado símbolo a una determinada respuesta. Pero considerada desde un punto de vista lógico, parecería ser el resultado de la categorización, puesto que el hecho de que la respuesta se ubique en una determinada categoría implica que se le atribuirá un determinado símbolo.

Probablemente el proceso cronológico real no haga admisible esta distinción en la mente del codificador, tanto para las preguntas llamadas erróneamente precodificadas (aquellas cuyas categorías de respuestas ya tienen un número de código asignado, impreso en la ficha de encuesta) como, sobre todo, para aquellas en las que el código es un documento independiente de las fichas de encuesta. Al buscar en el código a qué categoría corresponde la respuesta en cuestión, se encontrará, simultáneamente, el símbolo correspondiente (número); asimismo, al asignar el número de código (símbolo) se estará codificando y categorizando a la vez.

Tomemos como ejemplo de lo que estamos diciendo una investigación en la cual uno

de los datos a relevar sea el estado civil de los entrevistados. En la ficha de encuesta la pregunta podría estar diseñada de la siguiente manera:

Estado civil

Soltero	☐
Casado	☐
Divorciado	☐
Viudo	☐
Unión libre	☐
Sin dato	☐

Supongamos que una de las respuestas recibidas es “casado”. A los efectos de relevar el dato, puede marcarse en el cuadro correspondiente con una X, una raya inclinada, un círculo, o cualquier otro símbolo. Hasta el momento, lo único que hemos hecho es recoger el dato, pero aún no se ha planteado la etapa de codificarlo. Para hacerlo, en este caso, debemos recurrir a un elemento auxiliar, *el código* (al que nos referiremos más adelante), en el que encontraremos el número que corresponde a la respuesta “casado”. Ubicado el número, para completar la codificación podemos usar dos vías: escribir el número correspondiente junto a la respuesta en la propia ficha de encuesta, o recoger el número correspondiente a la respuesta en una hoja aparte (hoja de codificación).

Observamos que, además de codificar la respuesta “casado” con el número correspondiente, estamos colocando la respuesta, a los efectos de su elaboración (tabulación, etc.), dentro de una categoría. Dicha categoría comprende todas aquellas respuestas a las que corresponde asignar el mismo número. Vemos, entonces, el sentido que tenía nuestra afirmación: aunque se realicen simultáneamente, codificación y categorización implican dos etapas de la elaboración de los datos que deben analizarse por separado.

Si nos atenemos al ejemplo mencionado, es posible decir que la categorización de los datos, no ya a los efectos de su elaboración mecánica, sino simplemente de su ubicación en una etapa determinada, se realiza antes que la codificación. El solo hecho de marcar el cuadrado correspondiente al relevar el dato ubica la respuesta del sujeto dentro de las categorías generales de “soltero”, “casado”, etc., con independencia del número de código que le corresponda. Por cierto que esta posibilidad se da con referencia a cierto tipo de preguntas denominadas “cerradas”, pues admiten un número limitado de posibilidades de respuestas (en nuestro ejemplo 6 posibilidades).

Avancemos un poco más en nuestro ejemplo del estado civil, e imaginemos la misma pregunta diseñada de una manera distinta:

Estado civil

Soltero	☐ 1
Casado	☐ 2
Divorciado	☐ 3
Viudo	☐ 4
Unión libre	☐ 5
Sin dato	☐ 6

La diferencia con respecto al diseño anterior consiste en que se ha colocado un número al lado de cada casillero correspondiente a cada respuesta; el número indica el símbolo en que la codificación corresponde a cada respuesta. De modo que, al relevar el dato, se está al mismo tiempo categorizándolo y codificándolo. Para este tipo de preguntas que hemos llamado “cerradas”, el recurso de colocar en la propia ficha de encuesta el número de código tiene una finalidad eminentemente práctica: ahorra esfuerzo, evitando la duplicación innecesaria de tareas. Cuando directamente el encuestador, por el solo hecho de relevar el dato, codifica, libera a los codificadores de buscar cuál es el número que corresponde a cada respuesta de este tipo, con lo cual pueden dedicarse a la codificación de aquellas respuestas que no admiten este tipo de codificación simultánea al relevamiento.

Corresponde una última puntualización relativa a la necesidad de la codificación para poder elaborar los datos. A pesar de que prácticamente es empleada en todas las investigaciones sociológicas, no es una técnica necesaria para la elaboración de los datos desde un punto de vista estrictamente lógico. Es decir que su empleo no se impone por la naturaleza de los datos o por razones de orden metodológico. Puede concebirse una investigación en la que los resultados sean elaborados manualmente, es decir, que se revise ficha de encuesta por ficha de encuesta y se vaya categorizando en una hoja aparte las respuestas obtenidas. En el ejemplo del estado civil, tomaríamos una hoja en la que anotaríamos las 6 posibilidades de respuesta y posteriormente iríamos revisando cada ficha y anotando (con una tilde, etc.), la respuesta en esa hoja. Si tuviéramos que anotar las respuestas contenidas en 100 fichas, al final de nuestro trabajo tendríamos, en la hoja de elaboración:

Estado civil

	Total
Soltero III III III III III III	30
Casado III III III III III III III	33
Divorciado III III III	14
Viudo III III	9
Unión libre III III	10
Sin dato III	4
Total	100

De esta manera habríamos obviado el pasaje por la codificación, yendo directamente a contar la cantidad de casos que caen en cada una de las categorías de respuestas.

Lo que sí es insoslayable en toda investigación es la categorización de los datos, puesto que, según sea la frecuencia de respuestas que entren en cada categoría, la interpretación de los resultados de la investigación será distinta. Por sobre todas las cosas, las categorías en que se dividen las preguntas sirven para ubicar las respuestas en clases distintas y hacer posible, de esa manera, su ulterior interpretación.

Como se ha dicho, la codificación no es necesaria en un plano lógico, pero su empleo es corriente en las investigaciones sociológicas, fundamentalmente por dos razones. En primer término, porque el volumen de los datos relevados hacen necesaria su utilización. La elaboración de los datos, incluso en el caso de la elaboración manual, sin el auxilio de la codificación, es muy lenta y costosa, por encima de un número muy reducido de fichas de encuesta. La segunda razón radica en lo siguiente: cuando el volumen de los datos es muy grande, se impone su elaboración por medios mecánicos (con el sistema de tarjetas perforadas) y para poder utilizarlos la técnica de la codificación es imprescindible. Por otra parte, los símbolos numéricos en codificación se emplean teniendo en cuenta las limitaciones que imponen los métodos mecánicos de elaboración.

Distintas clases de preguntas

Conviene mencionar este punto antes de proseguir, puesto que la codificación se efectúa sobre la base de un código y éste es construido de manera distinta según el tipo de preguntas que integren el cuestionario. Pueden distinguirse, *grosso modo*, dos clases de preguntas fundamentales: preguntas abiertas y preguntas cerradas, las cuales, a su vez, pueden ser preguntas de hecho (o sobre hechos) y preguntas de opinión (o sobre

opiniones).

Preguntas abiertas se llama a aquellas en que el entrevistado puede contestar libremente; por tanto, el número de respuestas que puede darse resulta teóricamente ilimitado. Como ejemplo consignamos la siguiente pregunta: “¿Qué opina usted acerca de los partidos políticos?”

Por el contrario, las preguntas cerradas obligan al entrevistado a optar por una serie de posibles respuestas previamente consignadas en las mismas fichas de encuesta. Es el caso de nuestro ejemplo sobre el estado civil. La persona entrevistada debe, necesariamente, señalar una de las cinco posibilidades establecidas en la ficha de encuesta.

En cuanto a la categoría de preguntas relativas a hechos, tienen como finalidad indagar sobre situaciones referentes a las personas: estado civil, ocupación, edad, instrucción, ingresos, etc. El interés del encuestador está en que el entrevistado responda lo más verídicamente posible acerca de cualquiera de los rubros que en el cuestionario se mencionen, sin que interese lo que la persona opine al respecto.

Pregunta típica de opinión es la que mencionamos al ejemplificar el tipo de preguntas abiertas: “¿Qué opina usted acerca de los partidos políticos?” Como se comprende, no interesa saber lo que realmente son los partidos políticos, sino la opinión de los encuestados sobre los mismos. Las preguntas de opinión tienen como finalidad indagar acerca de las opiniones del grupo humano a entrevistar, sobre el o los tópicos que abarque la encuesta.

Tomando como base estas distinciones, podemos decir que, en general, las preguntas abiertas son de opinión y las preguntas cerradas son sobre hechos. Y decimos en general porque hay ciertas preguntas de hecho, como en el caso de las relativas a ocupación, que es prácticamente imposible diseñarlas como preguntas cerradas (si bien, técnicamente, deba admitirse).

De la misma manera, aunque las preguntas de opinión se diseñan generalmente abiertas (lo que se busca, justamente, es el máximo de espontaneidad en las respuestas), puede admitirse, y de hecho sucede, que una pregunta de opinión sea cerrada; por ejemplo, si se desea indagar sobre un tema cualquiera la existencia de estereotipos en la opinión pública.

Los problemas que se plantean en la confección de los códigos y en la técnica de la codificación están determinados, en parte, por el tipo de pregunta de que se trate. Con base en lo dicho anteriormente, estamos ahora en condiciones de hacer un breve comentario sobre la noción de código y su confección.

El código

En términos muy generales, puede decirse que código es el conjunto de símbolos correspondientes a cada una de las categorías en que se dividen las respuestas obtenidas en un cuestionario (ficha de encuesta, propiamente).

Se desprende que lo que interesa particularmente a la codificación son los símbolos

—en nuestro caso numéricos— atribuidos a cada categoría, más que las categorías propiamente dichas; éstas importarán, de modo fundamental, en las etapas posteriores de la elaboración de los datos. Sin embargo, debe tenerse presente que el código se confecciona sobre la base de la división de las respuestas obtenidas en categorías o clases. Por ello, para el establecimiento de las categorías se nos presentan todos los problemas —lógicos y prácticos— de la realización de una clasificación.

No nos detendremos a considerar los aspectos lógicos del problema; veremos sólo algunas de las indicaciones que señala la bibliografía sobre cuestiones más bien prácticas de la clasificación. Pero queremos destacar que los investigadores deben conocer en profundidad los problemas —tanto lógicos como prácticos— referentes a la clasificación, porque los errores cometidos en esta etapa de la investigación pueden comprometer la correcta verificación de la o las hipótesis.

Nomenclatura

Para Chevy,⁶ con el título de “Establecimiento de una nomenclatura”, ésta es “la lista de modalidades del carácter en cuestión que se pretende distinguir”. Es una clasificación de un ítem de cierta importancia, como, por ejemplo, las regiones geográficas, ocupación, industrias, causas de deceso, etc. El hecho de que una misma clasificación de un mismo ítem sea empleada en varias investigaciones permite una mejor comparación de los resultados, sobre todo para el tipo de encuesta al que se refiere Chevy, la encuesta estadística.

También podemos decir que nos hallamos frente a una nomenclatura cuando las respuestas a una pregunta deben dividirse en un número muy grande de rubros; sólo son consideraciones de cantidad de rubros e importancia del ítem —ambos aspectos ligados— las que, frente a casos concretos, nos permiten afirmar de un modo empírico la existencia de una nomenclatura.

Cuando nos hallamos ante una clasificación de oficios y profesiones que incluye centenares de rubros, diremos que estamos frente a una nomenclatura; no podríamos decir lo mismo ante una clasificación de sexo o estado civil; dos rubros en el primer caso, cinco o seis en el segundo.

Las reglas que Chevy señala como necesarias para el establecimiento de una buena nomenclatura pueden resumirse en: ser completa; ser lo suficientemente clara y precisa; no tener un carácter demasiado teórico, y adaptarse a las necesidades más generales en lo posible: “Decir que una nomenclatura debe ser completa no significa que debe ser complicada y detallada, sino al contrario. Conviene que sea lo suficientemente detallada como para proveer las clases de que se tenga necesidad; pero sería molesto que lo fuera en exceso”.

La necesidad de claridad y precisión surge de que “toda unidad estadística debe poder ser clasificada sin dudar en un rubro de la nomenclatura y sólo en uno”. “Para establecer una nomenclatura correcta que no presente ninguna ambigüedad, es necesario, previamente, elegir un criterio único de disociación de las modalidades que se requiere

distinguir.”

Puesto que la nomenclatura está destinada a clasificar las declaraciones relativas a los hechos, cuando se establece “no conviene conferirle un carácter demasiado teórico; no hacer una construcción del espíritu, satisfactoria desde un punto de vista lógico, pero que peligra no corresponder suficientemente con la realidad”.

Esta tercera regla merece, a nuestro juicio, dos observaciones que eviten confusiones. Primero, podemos aceptarla en la medida en que el autor está hablando de las encuestas estadísticas (o macrosociales) en las que importa el nivel descriptivo más que el explicativo. En segundo lugar, ese argumento es inaceptable para una investigación explicativa. Si no podemos efectuar una clasificación correcta desde el punto de vista lógico, es como consecuencia de nuestra ignorancia acerca de los hechos, no porque de la aplicación estricta de las reglas lógicas resulte un apartamiento de la realidad.

El propio Chevy se encarga de corroborar —en cierto sentido— nuestras observaciones: “Una nomenclatura estadística no tiene necesidad de constituir un instrumento de gran precisión, si debe servir para clasificar declaraciones, la mayor parte de las cuales serán relativamente vagas, lo que, desdichadamente, es el caso más general”.

Por su parte, Selltitz considera que deben seguirse tres reglas fundamentales para el establecimiento de las categorías: deben derivar de un principio simple de clasificación, es decir, deben ser exhaustivas, tiene que ser posible colocar cada respuesta en alguna de las categorías en que se halla dividida la pregunta, y las diversas categorías correspondientes a la misma pregunta deben ser mutuamente excluyentes, de manera que no sea posible colocar una respuesta dada en más de una categoría: “La calidad de la categorización de los datos depende enteramente de la validez de las categorías empleadas. Las categorías deben ser bien definidas desde un punto de vista conceptual y deben ser relevantes a los fines de la investigación”.

Antes de entrar en el tema de la confección del código parece pertinente dar una formulación un poco más precisa de cinco de los conceptos que se han manejado hasta ahora: clasificación, código, categorías, categorización y codificación.

Clasificación. Es la operación que permite ordenar —según ciertas reglas, de las cuales ya vimos algunas— las respuestas que se presume darán o dieron los entrevistados (dependiendo de que la clasificación se realice antes o después del relevamiento) al cuestionario que les fue sometido. Esas respuestas, por lo tanto, se ubican en clases distintas.

Código. Es el conjunto de símbolos que se atribuye a cada una de las diferentes clases en que se dividen las respuestas a una pregunta.

Categorías. Son, precisamente, las clases en que vimos que se dividen las respuestas a una pregunta. Conceptualmente, parece que no puede establecerse ninguna distinción entre clase y categoría. Los autores consultados emplean indistintamente ambos conceptos, aunque el término categoría es, en general, más empleado. Podría establecerse una sutil diferencia entre clase y categoría si sostenemos que las clases pasan a ser categorías cuando se les asigna un número de código, o cuando esas clases

son recogidas en un documento especial con vistas a la elaboración de los datos; pero sería peligroso sostener esta distinción como evidente.

Categorización. Consiste en ubicar una respuesta concreta en una de las clases o categorías.

Codificación. Supone la atribución del símbolo —extraído del código— que corresponde a la respuesta en cuestión.

En el marco de estas cuatro operaciones —clasificación, construcción del código, categorización y codificación— están inscritos los problemas referentes a la codificación.

Confección del código

Para Duverger,⁷ la confección del código “no ofrece ninguna dificultad en las preguntas cerradas, en las que, por definición, las respuestas son limitadas y fáciles de distinguir... También es muy fácil en las preguntas de hecho: sexo, edad, residencia, partido, religión. Por el contrario, es muy difícil en las preguntas abiertas de respuesta totalmente libre”.

Conviene acotar que cuando Duverger hace referencia al código le da, justamente, el sentido con que nosotros lo definimos: el conjunto de símbolos correspondientes a cada una de las categorías y no las categorías propiamente dichas. Es en este sentido como puede aceptarse su afirmación de que la elaboración del código no ofrece dificultades en las preguntas cerradas, sin caer en el error de confundir esto con categorización. Aun tratándose de preguntas cerradas, el previo establecimiento de las categorías presenta dificultades que hay que resolver bien para no malograr los resultados de una investigación.

A los efectos del diseño y confección del código, interesa primordialmente la división en preguntas abiertas y cerradas, más que la división entre preguntas de hecho y de opinión; razón por la cual al hablar de preguntas cerradas englobaremos en esta denominación tanto las de hecho como las de opinión. Más adelante veremos que esta otra distinción nos interesa cuando hablemos de preguntas abiertas.

Como ya mencionamos, las preguntas cerradas son aquellas en las que todas las posibilidades de respuesta del entrevistado están taxativamente enumeradas de antemano.

Tomemos, por ejemplo, la siguiente pregunta: “¿Está usted de acuerdo con la actual política del gobierno?” Las posibilidades de respuesta son infinitas, pero por razones que tienen relación con el objetivo de la investigación nos interesa solamente saber a grandes rasgos la opinión de los entrevistados y, en este sentido, “cerramos” la pregunta en las siguientes posibilidades:

¿Está usted de acuerdo con la actual política del gobierno?

- | | |
|--------------------------|--------------------------|
| Completamente de acuerdo | ☐ |
| Medianamente de acuerdo | ☐ |
| Totalmente en desacuerdo | ☐ |
| No desea contestar | ☐ |
| Sin dato | ☐ |

En este caso, de igual modo que en el ejemplo dado sobre el estado civil, podemos colocar un número en el costado exterior derecho de cada casillero, de tal forma que la pregunta quedaría confeccionada de la siguiente manera:

¿Está usted de acuerdo con la actual política del gobierno?

- | | |
|--------------------------|----------------------------|
| Completamente de acuerdo | ☐ 1 |
| Medianamente de acuerdo | ☐ 2 |
| Totalmente en desacuerdo | ☐ 3 |
| No desea contestar | ☐ 4 |
| Sin dato | ☐ 5 |

De la misma manera es posible proceder con todas las preguntas cerradas que deban ubicarse en una ficha de encuesta; es decir que los números correspondientes a cada categoría de respuestas a cada pregunta están impresos en la propia ficha.

De modo que podemos decir que el código de esta clase de preguntas está inserto en la propia ficha de encuesta, en vez de formar un instrumento separado, como veremos que sucede con las otras clases de preguntas. Si todas las preguntas de una encuesta fueran cerradas, la ficha y el código constituirían un solo y único documento.

Pasemos ahora a considerar los problemas que se plantean en la confección del código cuando se trata de preguntas abiertas, considerando en primer lugar las referidas a cuestiones de hecho y tomando como ejemplo típico el de la ocupación.

Si bien teóricamente es posible, en los hechos resulta imposible colocar en una ficha de encuesta todas las posibles ocupaciones de los sujetos a entrevistar; sobre todo si se tiene en cuenta que la ocupación es una variable de gran importancia y que se impone relevarla con cierta minuciosidad. Por otro lado, cerrar la pregunta sobre ocupación complicaría enormemente el diseño de la ficha de encuesta, no compensando los riesgos que se correrían en este caso frente al tiempo que lleva la codificación posterior de las respuestas.

La solución que se adopta, en consecuencia, es colocar en la ficha cada uno de los ítems de la variable, con un determinado espacio en blanco para que los encuestadores escriban el dato que corresponda. El diseño puede ser como sigue:

Ocupación

Trabajo
Dependencia
Sector
Rama
Lugar

Una vez relevados los datos requeridos, en el casillero (o casilleros, según el código implique la colocación de uno, dos o más números), el codificador coloca el o los números correspondientes, según el código, junto a cada casillero.

Esta operación implica que, con anterioridad al relevamiento o, por lo menos, con anterioridad a la codificación, se ha confeccionado un código de las ocupaciones. Este código consiste en un instrumento separado y distinto de la ficha de encuesta, en el que las ocupaciones o los ítems, en su caso, son divididos en diversas categorías, a las que se les asigna uno o varios números, según corresponda.

No se nos escapa, y nos parece importante señalarlo con cierta extensión por la frecuencia con que se investiga en el tema, que existe una gran dificultad para categorizar las ocupaciones. Esta cuestión no ha sido hasta ahora resuelta de manera uniforme; en cada país, es más, en cada investigación, es frecuente encontrar categorizaciones diferentes sobre similares ocupaciones.

Hay dos tipos de problemas fundamentales que se plantean en el caso de la categorización de ocupaciones.

Por un lado, la construcción del código es muy difícil por cuanto la elaboración de los datos por el sistema de tarjetas perforadas impone limitaciones a la utilización de los símbolos numéricos; es necesario hacer la combinación de números de código en vista a una elaboración eficaz, que evite demoras y confusiones. Por otro lado, la codificación en sí presenta dificultades; no siempre es fácil ubicar con precisión una ocupación en determinada categoría, sobre todo cuando se trata de casos límite, los cuales son más frecuentes de lo que podría suponerse.

Otra cuestión en relación con la codificación de las ocupaciones, que además tiene validez general, es la forma en que ella se realice. Si se trata de preguntas de hecho, la codificación de respuestas a preguntas abiertas puede hacerse directamente en la ficha de encuesta, si es que se han previsto los casilleros a ese efecto. Pero también la codificación puede realizarse (y es lo corriente en el caso de las respuestas a preguntas abiertas de opinión) en una hoja separada, comúnmente denominada hoja de codificación.

En relación con la confección del código para las preguntas abiertas de opinión, Duverger⁸ resume con mucha claridad el procedimiento a seguir: “La elaboración del código se realiza en varias etapas. En primer lugar, se establece un ‘precódigo’, fundado en un análisis hipotético de los tipos de respuestas posibles. El precódigo se aplica después a una muestra de entrevistas, lo que permite rectificar experimentalmente sus categorías y establecer el código propiamente dicho, que, por lo general, es sometido

nuevamente a una verificación experimental llamada ‘codificación de prueba’, en la que varios investigadores codifican separadamente una selección de entrevistas, comprobándose así la fidelidad del código”.

Hyman,⁹ en la sección relativa a las instrucciones para codificar las respuestas a preguntas abiertas, señala que la forma de desarrollar el código es a través de “la lectura de una muestra de respuestas”, lo que, si bien mucho menos explícito, coincide con lo dicho por Duverger.

<i>Núm. Pregunta-Variable</i>	<i>Columna</i>	<i>Código</i>	<i>Alternativas</i>
	1-3		Número de la entrevista
1 Edad	4-5		Las dos últimas cifras del código corresponden a la fecha de nacimiento. Alguien nacido el año 1936 tiene entonces el número de código 36; una persona nacida en el año 1897 tiene código 97, etc.
Sexo	6	1 2	Masculino Femenino
2 Estado Civil	7	1 2 3	Casado Soltero Ha estado casado (divorciado, viudo)
3 Grado de Urbanización	8	1 2 3 4 5 6	Capital Ciudades sobre 250 000 Ciudades urbanas con menos de 250 000 pero más de 110 000 Ciudades con menos de 110 000 pero más de 10 000 Comunidades con menos de 10 000 no rural Rural
4 Ingreso Familiar 1968 al mes (jefe de familia)	9	1 2 3 4 5 6 7 8 9	Sin ingreso Menos de E° 200 Más de E° 200 y menos de E° 400 Más de E° 400 y menos de E° 800 Más de E° 800 y menos de E° 1 200 Más de E° 1 200 y menos de E° 2 000 Más de E° 2 000 y menos de E° 2 800 Más de E° 2 800 y menos de E° 4 400 Más de E° 4 400
5 En qué sector industrial trabaja el jefe de hogar	10	1 2 3 4 5 6 7 8 9	Agricultura, silvicultura Construcción de casas Industria y artesanía Transporte Comercio Empleado público administrativo Profesiones liberales Otras ocupaciones Sin profesión, por ej., dueña de casa, jubilado, estudiante, etc.

En cuanto a la distinción que puede hacerse, desde el punto de vista de la confección del código, entre las preguntas abiertas de hecho y las preguntas abiertas de opinión, puede decirse que radica, en general, en que las preguntas sobre hechos admiten que el código se confeccione antes del relevamiento; mientras que, en el caso de las relativas a opinión, el código es redactado con posterioridad al relevamiento, sobre la base del examen de una muestra de las entrevistas recogidas. Esta última técnica se impone por la propia índole de la pregunta; no es posible prever qué es lo que contestarán los entrevistados, o mejor, no es posible preverlo en un grado tal que permita confeccionar de antemano las categorías.

Como consecuencia para la etapa de codificación, las categorías así establecidas tienen una imprecisión mayor que en el caso de preguntas sobre hechos, y se da más fácilmente la posibilidad de que existan respuestas que no encuentren cabida en ninguna de las categorías establecidas. Una dificultad de esta naturaleza impone el esfuerzo de rehacer el conjunto de categorías en que ha sido dividido un ítem, en vista de que dicha división ha resultado inadecuada para ordenar el total de respuestas.

Ejemplos de códigos

A continuación anotamos algunos ejemplos de diferentes tipos de códigos obtenidos en diferentes investigaciones sociológicas.

Normalmente, en el código se anota el número de la pregunta; en algunas ocasiones, el nombre de la variable, la columna en el espacial IBM, el código, es decir, la posición, dentro de las columnas y alternativas utilizadas. (*Cuadro en la página anterior.*)

Para facilitar el procesamiento de datos, conviene trabajar con un código general para las siguientes categorías:

O = Según las instrucciones, esta pregunta no se aplica.

X = El entrevistador no ha puesto la pregunta porque ha encontrado impropio hacerlo, por alguna razón. (Indique por qué.)

Y = La pregunta fue hecha al entrevistado el cual no quiso responder o no sabía la respuesta o ha respondido de manera incomprensible.

Si el entrevistador, por alguna razón, quiere dividir los entrevistados que han obtenido una *Y* en algunas preguntas, en un grupo que ha rechazado responder, corresponde a *NA* = *no answer* y los que quieren contestar, pero no saben, *DK* = *Don't know*. Se puede marcar *Y* en el cuestionario con una explicación del tipo de *Y*.

Este código aparece, entonces, para cada pregunta pero no está escrito en el código solamente. Por eso dejamos como alternativas en la columna las posiciones 1, 2..., 8, 9, libres para las alternativas. Si el investigador está utilizando algunos tipos de computadoras, tendrá que cambiar el código porque la máquina no opera con las posiciones *X* y *Y*.

El código puede ocupar también dos columnas, si el número de alternativas de una pregunta son muchas:

Núm.	Pregunta-Variable	Columna	Código	Alternativas
				<i>Razones para el peligro (1)</i>
6	En su opinión, ¿qué es lo que probablemente sucederá en la situación mundial del próximo año o en los próximos dos años?	11-12	11	Individualismo, líderes, personas
			12	Líderes soviéticos (como individuos)
			13	Líderes norteamericanos (como individuos)
			14	Otras personas especialmente mencionadas
			21	<i>Desarrollo político internacional entre países (2)</i>
			22	La estructura rusa (poder político, etc.)
			23	La estructura política norteamericana (poder político, etc.)
			24	Otras condiciones especiales de nacionalismo político (la división de Alemania, etc.)
			31	<i>Diferencias económicas y sociales (general)</i>
			32	Competencia de utilidades, la dominación económica de los grandes poderes, competencia de mercados, etc.
			33	La gente menos favorecida luchando por obtener un estándar más alto
			34	Tendencia a los cambios en la estructura de clases
			38	Otras diferencias especiales, económicas y sociales
			41	<i>Sistemas políticos o ideologías (general)</i>
			42	Comunismo
			43	Capitalismo
			44	Influencias nacionalistas
			48	Otras influencias especiales, políticas o ideológicas
			51	<i>Étnicas, diferencias culturales — diferencias entre gente (general)</i>
			52	La consistencia de las culturas occidental y oriental
			53	Desacuerdo entre gente, desconfianza, etc.
			54	Tensiones entre razas
			58	Otras diferencias especiales, culturales y étnicas
			61	<i>Poderees históricos (general)</i>

<i>Núm.</i>	<i>Pregunta-Variable</i>	<i>Columna</i>	<i>Código</i>	<i>Alternativas</i>
			62	El desarrollo científico
			63	El desarrollo social (sociedades más complejas, problemas demográficos, migraciones, etc.)
			68	Otras tendencias históricas especiales
			71	<i>Moral-condiciones étnicas (general)</i>
			72	Moral. Situación ética en el Hemisferio Occidental
			73	Moral. Situación ética en el Hemisferio Oriental
			78	Otras condiciones específicas de moral y ética
			81	<i>Insuficiencia de las organizaciones internacionales (general)</i>

Aquí se puede pensar también en una *ordenación* dentro del código, como en este ejemplo:

<i>Núm.</i>	<i>Pregunta-Variable</i>	<i>Columna</i>	<i>Código</i>	<i>Alternativas</i>
	Clase Social	13	1	Pudiente (acomodada)
			2	Clase media
			3	Obreros (trabajadores) y sus similares
7	Ocupación	13-14	11	Dueños de fundo y arrendatarios de importancia
			12	Industriales
			13	Empresarios mayoristas
			14	Altos empleados de empresas
			15	Altos empleados públicos
			16	Profesiones liberales, ej.: abogados, arquitectos, dentistas
			17	Propietarios de grandes hoteles y restaurantes
			18	Propietarios de edificios de departamentos, casas, etc.
			21	Agricultores (predio mediano), incluyendo parientes cercanos que trabajen en el fundo
			22	Empresarios artesanales
			23	Dueños de tienda, comerciante (pequeños)

<i>Núm.</i>	<i>Pregunta-Variable</i>	<i>Columna</i>	<i>Código</i>	<i>Alternativas</i>
			24	Funcionarios de posiciones intermedias en empresas
			25	Funcionarios públicos de categoría media
			26	Profesiones liberales y medias
			27	Empleados posición media en hoteles y restaurantes
			28	Servicios domésticos (medias)
			31	Trabajadores agrícolas
			32	Artesanales (baja categoría)
			33	Trabajadores (actividades comerciales)
			34	Trabajadores (industria y minería)
			35	Funcionarios públicos de bajas posiciones
			36	Profesiones liberales (inferiores)
			37	Lavanderas, etc.

La primera cifra (correspondiente a la columna 13) da la ordenación de la ocupación, con respecto a clase (baja, media, alta) (tricotomización). La segunda cifra da el sector de trabajo.

Entonces, el dueño de un fundo tiene el código 11, en columna 13-14; el subordinado 21 y el trabajador agrícola 31.

Otra manera de utilizar el espacio sería la siguiente:

<i>Núm.</i>	<i>Pregunta-Variable</i>	<i>Columna</i>	<i>Código</i>	<i>Alternativas</i>
8	¿Es usted miembro de los siguientes clubes o asociaciones? En caso afirmativo, ¿cuál es su posición dentro del club? Club de deportes	15-17	15:1	Perforado- <i>Ip</i> es miembro de la directiva
			16:1	" - <i>Ip</i> tiene una misión especial dentro del club
			17:1	" - <i>Ip</i> es un socio corriente
			15-17:1	Sin perforar- <i>Ip</i> no es miembro de ningún club de deportes
9	Agrupación política	15-17	15:4	Perforado- <i>Ip</i> es miembro de la directiva
			16:4	" - <i>Ip</i> tiene una misión

<i>Núm.</i>	<i>Pregunta-Variable</i>	<i>Columna</i>	<i>Código</i>	<i>Alternativas</i>
				especial dentro de la agrupación
			17:4	Perforado- <i>Ip</i> es miembro corriente
			15-17:4	Sin perforar- <i>Ip</i> no es miembro de ninguna agrupación política
10	Liga de templanza	15-17	15:5	Perforado- <i>Ip</i> es miembro de la directiva
			16:5	" - <i>Ip</i> tiene una misión especial dentro de la liga
			17:5	" - <i>Ip</i> es miembro corriente
			15-17:5	Sin perforar- <i>Ip</i> no es miembro de ninguna liga de templanza
11	Sociedad disidente	15-17	15:6	Perforado
			16:6	" Mismas alternativas
			17:6	"
			15-17:6	Sin perforar- <i>Ip</i> no es miembro de ninguna sociedad disidente
12	Sociedad religiosa (Iglesia del Estado)	15-17	15:7	Perforado
			16:7	" Las mismas alternativas
			17:7	"
			15-17:7	Sin perforar- <i>Ip</i> no es miembro de ninguna sociedad religiosa
13	Club musical o de canto	15-17	15:8	Perforado- <i>Ip</i> es miembro de la directiva
			16:8	" - <i>Ip</i> tiene una misión especial dentro del club
			17:8	" - <i>Ip</i> es miembro corriente
14	Cruz Roja	15-17	15:9	Perforado
			16:9	" Las mismas alternativas
			17:9	"
			15-17:9	Sin perforar-No pertenece a la Cruz Roja

Finalmente presentamos una solución aún un poco pesada con respecto al problema de múltiples alternativas.

Supongamos que necesitamos saber si el individuo entrevistado tiene algún conocimiento de idiomas extranjeros. Para facilitar la presentación, tomamos solamente la alternativa de si sabe leer los idiomas. Entonces, existe la posibilidad de que sabe leer: ninguno, uno, dos o más idiomas. Si sabe leer tres idiomas, pueden ser tres diferentes idiomas que maneja otro individuo, etc.

<i>Núm.</i>	<i>Pregunta-Variable</i>	<i>Columna</i>	<i>Código</i>	<i>Alternativas</i>
15	¿Qué idiomas sabe usted leer?	18-19	01	Inglés
			02	Francés
			04	Alemán
			10	Italiano
			20	Portugués
			40	Otros

Este tipo de código está construido en la idea binaria de resolver el problema. Basta, entonces, con sumar las cifras indicadas para cada idioma y el código total va a ser el siguiente:

Núm.	Pregunta-Variable	Columna	Código	Alternativas
			01	Inglés
			02	Francés
			03	Inglés y francés
			04	Alemán
			05	Inglés y alemán
			06	Francés y alemán
			07	Inglés, francés y alemán
			10	Italiano
			11	Italiano e inglés
			12	Italiano y francés
			13	Italiano, inglés y francés
			14	Italiano y alemán
			15	Italiano, inglés y alemán
			16	Italiano, francés y alemán
			17	Italiano, inglés, francés y alemán
			20	Portugués
			21	Portugués e inglés
			22	" y francés
			23	" inglés y francés
			24	" y alemán
			25	" inglés y alemán
			26	" francés y alemán
			27	" inglés, francés y alemán
			30	Italiano y portugués
			31	Italiano, portugués e inglés, etc.

Procedimiento de codificación

Para Duverger,¹⁰ la codificación propiamente dicha consiste en “introducir cada respuesta en una de las diversas categorías definidas por el código”. Nosotros preferimos decir que consiste en atribuir a cada respuesta el símbolo numérico correspondiente a la categoría “definida del código” en que cae la respuesta. Asimismo, recordamos la mención de Selltitz, cuando decía que no es un procedimiento automático, sino que envuelve un juicio por parte del codificador.

Por otra parte, esa tarea puede ser hecha no sólo por los codificadores, sino por los propios encuestadores en el momento del relevamiento, lo que nos obliga a dividir la discusión sobre el procedimiento de codificar según esta operación sea hecha por unos u otros. Lógicamente, nos importa describir con más detalle cómo se realiza la labor de los codificadores, en razón sobre todo del siguiente hecho: los datos de la ficha que se dejan para codificar luego del relevamiento son los más complejos y requieren un entrenamiento técnico adecuado por parte de la persona que va a codificarlos; pese a que los encuestadores también codifican, la clase y nivel de su adiestramiento es distinto del

de los codificadores y tiene relación directa con la técnica de la entrevista, problema que escapa al ámbito de este trabajo.

Incluso podemos señalar un aspecto práctico de este dualismo en cuanto a las personas que codifican; en el lenguaje de los centros de investigación que conocemos en Uruguay, se llama codificación a la etapa que se efectúa luego del relevamiento, realizada por codificadores, aunque se codifica también durante el relevamiento. Esta manera de dominar esa segunda etapa es un buen reflejo de su mayor importancia frente a la codificación realizada por los encuestadores.

Codificación que realiza el encuestador

Según Selltitz,¹¹ las ventajas que se derivan de esta forma de codificación son las siguientes: el encuestador está en una situación tal que le permite captar con acierto la respuesta del entrevistado; en virtud de ello, tiene más información sobre la cual basar su juicio, con relación al codificador que trabaja sobre la base de datos escritos; y este tipo de codificación ahorra tiempo y trabajo. “No obstante estas ventajas [anotan los autores mencionados], la categorización [nosotros diríamos codificación] de datos complejos es generalmente efectuada por codificadores después que los datos han sido relevados.”

El procedimiento de codificación en la oficina

Llamamos así a la codificación que realizan los codificadores, propiamente, a los solos efectos de distinguirla de la que realizan los encuestadores en el momento de la entrevista.

El primer problema a considerar es el referente al entrenamiento de los codificadores. Por lógica, una correcta codificación de los datos obtenidos en el relevamiento tiene como supuesto previo el contar con un equipo de codificadores adecuadamente adiestrados. El periodo de preparación será más o menos largo, según la complejidad de los datos a codificar, y puede ser llevado a cabo antes, durante o después del relevamiento. Pero, como se comprende, es más prudente tener el elenco de codificadores entrenado de manera tal que pueda entrar en funciones inmediatamente después del relevamiento, evitando así pérdida de recursos económicos y de tiempo.

Siguiendo a Selltitz, mencionaremos las etapas que deben transitarse en el entrenamiento de los codificadores. Asimismo señalamos que en esta parte, así como en el resto de la exposición, seguiremos de cerca a los autores que han desarrollado con extensión estos puntos, puesto que consideramos útil un planteamiento eminentemente práctico del tema.

Las etapas fundamentales en el entrenamiento de codificadores serían las siguientes:

En primer término, el código o los códigos a emplear son explicados e ilustrados con ejemplos extraídos del material a ser codificado. Los codificadores, entonces, practican sobre una muestra de los datos recogidos. Los problemas que surgen son discutidos por los codificadores en grupo, con el supervisor, con la finalidad de desarrollar procedimientos y definiciones comunes.

Frecuentemente, como resultado de esta práctica de codificación, las categorías ya establecidas son revisadas para hacerlas más aplicables al material relevado y para poner en práctica los procedimientos y definiciones que han sido desarrollados durante la codificación preliminar.

En el periodo de práctica, cuando aparecen pocos problemas, los codificadores trabajan sobre una porción idéntica de los datos, sin consultar a los otros codificadores o al supervisor. La consistencia y precisión de la codificación es entonces computada para determinar si es factible comenzar a codificar formalmente.

Por su parte, Chevy,¹² bajo el título “Dificultades de la codificación”, nos dice: “La codificación es una operación que presenta dificultades extremadamente variables, según la naturaleza de las preguntas hechas, la calidad de las respuestas y la complejidad del código utilizado”.

A su juicio, las dificultades que pueden plantearse tienen, generalmente, una o más de las siguientes causas: el concepto que se posea del asunto u objeto a codificar: “para servirse eficazmente de un código de diplomas universitarios, es indispensable tener una idea bastante precisa sobre la organización de la enseñanza”. Otra dificultad la plantea la nomenclatura en sí misma. Sostiene que, por más cuidado que se tenga en la confección de las grandes nomenclaturas, éstas “no son jamás suficientemente precisas y completas para que su empleo no pueda dar lugar a dudas y a dificultades de interpretación”. En tercer término, la gran variedad de dificultades que pueden presentarse en las codificaciones, así como el carácter muy especial de alguna de ellas.

Esto último ha llevado a los servicios oficiales de estadística, que deben hacer codificaciones numerosas y variadas, a clasificarlas según su grado de dificultad y a especializar su personal, fundamentalmente, en dos aspectos: en presencia de cuestionarios que comportan numerosas preguntas en las cuales la codificación presenta dificultades desiguales, no se pide a la misma persona que codifique todas las preguntas... de tal modo que las codificaciones más simples sean confiadas al personal novato y las más difíciles al personal más capaz y mejor formado.

En segundo término, “es aconsejable que las codificaciones particularmente delicadas, como las causas de deceso o actividades económicas, sean siempre confiadas a las mismas personas especialmente formadas a este efecto y que adquieren de este modo una competencia que garantiza a la vez la calidad y la rapidez de su trabajo”.

Pasando ahora a los aspectos referidos concretamente al procedimiento de la codificación, mencionaremos a Hyman, quien detalla las instrucciones para los codificadores empleados en el National Opinion Research Center:

La regla fundamental en codificación es: sea cuidadoso. Si usted asigna un código equivocado a una respuesta o escribe sus números de modo tan ilegible que el *puncher* —persona que realiza la tarea de trasladar los datos codificados a tarjetas perforadas— los interpreta equivocadamente, los pone en una columna que no corresponde, los resultados del *survey* son distorsionados. De tal modo, nosotros le indicaremos las siguientes reglas generales a seguir: 1) Use lápiz rojo; 2) Escriba sus

números claramente; 3) No borre. Tache el material incorrecto.

A continuación se mencionan instrucciones para trabajar con dos tipos de preguntas: preguntas precodificadas, es decir, el tipo de preguntas cerradas que son directamente codificadas por el encuestador; y preguntas abiertas, vale decir, aquellas cuya codificación se realiza en la oficina.

Como se comprende, para el tipo de preguntas precodificadas la labor del codificador consiste en revisar si el dato ha sido bien relevado, y algunas de sus tareas, a título de ejemplo, son las siguientes, según Hyman:

Primera: comprobar que cada pregunta que debía hacerse ha sido realmente hecha y que el casillero correspondiente (con el número de código ya asignado) ha sido marcado.

Segunda: comprobar que una determinada pregunta no ha sido hecha si no correspondía efectuarla.

Tercera: en caso de que el encuestador haya marcado erróneamente el casillero que corresponde a una respuesta, el codificador debe enmendar la codificación.

Cuarta: si el encuestador pasa por alto una pregunta que debía ser hecha, el codificador debe colocar el símbolo correspondiente a “pregunta no hecha”, “sin dato”, etc. Tanto la denominación como el símbolo para codificar este tipo de pregunta “en blanco” varían según los investigadores, no existiendo un acuerdo general sobre la cuestión. De cualquier manera, conviene señalar que comúnmente se utiliza la letra *Y* para las preguntas no hechas, símbolo que corresponde al número 11 en las tarjetas perforadas tipo IBM, y la letra *X* para las preguntas acerca de las cuales no pudo obtenerse información (sin dato), símbolo que corresponde al número 12 de las tarjetas; para este último caso, se emplea también en forma habitual directamente el número 12, símbolo que es cada vez más general para codificar “sin dato”.

Quinta: si el encuestador omitió marcar el casillero correspondiente, pero escribió lo que el entrevistado dijo, el codificador debe leer los comentarios y, si éstos lo permiten, escribir el número pertinente; en caso contrario, debe codificarse o bien “no sabe” o bien “sin dato”.

Por último, si el encuestador ha marcado dos casilleros pertenecientes a la misma pregunta y no hay comentario alguno escrito sobre la cuestión, que permite saber qué número de código corresponde, debe codificarse “no sabe”.

El resto de la tarea de contralor de la codificación realizada por los encuestadores compete más directamente a los procedimientos generales de revisión a que es sometida una ficha de encuesta, a la que nos referiremos más adelante; en los hechos, escapa a la tarea propia de los codificadores.

Instrucciones para la codificación de preguntas abiertas

Como ya señaláramos, la codificación de preguntas abiertas —ya sean de hecho o de opinión— puede ser efectuada directamente en la ficha de encuesta o en una hoja u hojas separadas.

Siempre con base en el libro de Hyman,¹³ veremos algunos aspectos de la codificación de preguntas abiertas. Es importante, en este caso, familiarizarse no sólo con el código, sino además con las diferentes categorías en que se dividió cada pregunta, de modo tal que no haya dudas acerca de la codificación de las respuestas.

Siempre que corresponda hacer la pregunta y ésta haya sido contestada, debe efectuarse la codificación. En el caso de respuestas totalmente inadecuadas a la pregunta —lo que, de paso, revelaría una deficiencia del encuestador que relevó los datos—, no puede dejar de codificarse con el pretexto de que no se puede ubicar con precisión la respuesta de una categoría.

Deben codificarse las ideas y no las palabras. El codificador no puede pretender encontrar, en la definición y ejemplos de las categorías pertinentes a la respuesta a la pregunta, la frase o frases que correspondan exactamente con la respuesta del entrevistado.

Generalmente, se coloca una categoría denominada “otras respuestas” en cada pregunta (sobre todo en las preguntas abiertas de opinión) para prever la probabilidad de que una respuesta no pueda ser codificada dentro de las otras categorías establecidas; pero no deben colocarse en esta categoría todas aquellas respuestas vagas, ininteligibles, etc. Deben ubicarse las que signifiquen posibilidades nuevas de respuestas, no previstas en el elenco de categorías. La categoría “otras respuestas” cumple, fundamentalmente, la siguiente función: si el número de respuestas que entran en esta categoría es muy elevado será necesario crear nuevas categorías y revisar la división de categorías original.

Procedimientos administrativos. Algunos ejemplos

Además de las instrucciones y procedimientos de orden técnico en el trabajo de codificación, el mismo implica una organización y procedimientos administrativos: organización y control del trabajo de los codificadores, distribución de tareas, manejo y control desde el punto de vista del desplazamiento interno de las fichas de encuesta, etcétera.

Transcribiremos los aspectos sustanciales que menciona Hyman bajo el título “Procedimientos administrativos”:

Los cuestionarios son numerados a medida que son recibidos y agrupados en paquetes de 100 (cuestionarios) cada uno. Un paquete es la unidad con la cual usted trabajará en un tiempo dado. Como regla general, los codificadores deberán codificar sólo una página de un paquete determinado. De esta forma se familiarizan con el código de las preguntas de esa página más rápidamente que si tuvieran que codificar el total del cuestionario al mismo tiempo.

Hemos establecido varias normas referentes a las operaciones de codificación, las cuales le pedimos que observe:

1. Si tiene dudas acerca de qué número de código colocar en un cuestionario,

consulte al supervisor, no a sus compañeros codificadores.

2. Si desea consultar al supervisor un problema de codificación, espere hasta terminar un paquete sobre una página determinada.

3. Por favor, mantenga en un mínimo las discusiones durante el periodo de trabajo.

4. Es responsabilidad de los codificadores de la página 1 asegurarse de que el cuestionario ha sido correctamente procesado. La siguiente revisión debe ser hecha por aquellos que codifican la página 1: *a)* Ver que el cuestionario esté adecuadamente clasificado (*stapled*). *b)* Ver que cada cuestionario tenga asignado un número y que no haya números omitidos o duplicados. Cualquier deficiencia en la numeración debe ser informada de inmediato al supervisor, *c)* Reescriba cualquier número de un cuestionario que sea ilegible.

5. Los codificadores, en general, deben vigilar la existencia de páginas duplicadas u omitidas en el cuestionario.

6. Todo el material es considerado confidencial.

Procedimiento de revisión y control

Durante o después de la etapa de codificación, los propios codificadores u otras personas —más especializadas y con mayor dominio de la ficha de encuesta— son requeridos para revisar la codificación que se ha efectuado. En general, no se controlan todos los cuestionarios, lo que demandaría mucho tiempo, sino un cierto porcentaje de ellos: un tercio, un medio, etc. Más que la corrección de los errores cometidos, lo que interesa es conocer el porcentaje de los mismos en el conjunto de cuestionarios; a este porcentaje se le denomina “margen de error” (téngase presente que aquí hablamos de “margen de error” en la codificación de los datos). El margen de error que puede admitirse varía de una investigación a otra y se fundamenta en consideraciones metodológicas que no es pertinente discutir aquí. Interesa señalar que, en esta etapa, el cuestionario es sometido a lo que se denomina crítica de consistencia interna: se trata de verificar que todas las respuestas al cuestionario sean congruentes entre sí, vale decir, por ejemplo, que una persona de 24 años no aparezca luego como casada 15 años atrás. Para Chevy:¹⁴

La codificación de los cuestionarios debe ser objeto de una verificación consistente en hacer revisar el trabajo realizado por una empleada más calificada que la codificadora y capaz, por consiguiente, de rectificar los errores que esta última ha podido cometer. Pero es evidente que entendida de este modo la verificación de la codificación sería un trabajo muy oneroso, pues casi doblaría el costo de la codificación en sí misma. Como consecuencia, se limita a menudo, cuando se trata de trabajar con grandes masas de documentos:

Sea a la verificación de las codificaciones más importantes y más difíciles, sea a verificar nada más que una parte de los cuestionarios y a seguir la calidad del trabajo de cada codificadora merced a una verificación por sondeo de los documentos que le han sido confiados. Se asegura de este modo que el porcentaje de errores que comete

la codificadora no pase de una tasa que se ha fijado de antemano, y se procede a una verificación integral de los lotes de documentos para los cuales ella hubiera pasado esa tasa. Bien entendido, este procedimiento no permite rectificar todos los errores de codificación; tiene sin embargo la ventaja de ser económico y limitar el porcentaje de errores residuales a un nivel admisible.

Para evitar errores de codificación que introducen anomalías e incompatibilidades muy difíciles de descubrir y corregir más tarde, es muy recomendable proveer a las verificadoras y eventualmente a las codificadoras de tablas de incompatibilidades.

Señalamos, finalmente, que para la revisión y control se aconseja emplear una hoja aparte y separada de la de codificación, en la que el encargado de esta labor anota sus desacuerdos, las omisiones, errores de la codificación original, etcétera.

Para terminar nuestra discusión de los procedimientos de revisión y control, transcribiremos algunas de las consideraciones que menciona Selltiz,¹⁵ bajo el título de “Problemas de la confiabilidad en la codificación”:

Cada entrevista o cédula de observación debe ser revisada para:

1) Completarla. Todos los ítems deben ser llenados. Un claro a continuación de una pregunta en una cédula de entrevista puede significar “no sabe”, “rehúsa contestar”, o que la pregunta no es pertinente, que la pregunta fue omitida por error, etc. Para nuestros propósitos, es importante ser capaz de distinguir entre estos significados potenciales.

2) Hacerla legible. Si el codificador no puede descifrar lo que ha escrito el encuestador u observador, o las abreviaturas o símbolos empleados por el mismo, la codificación se torna imposible.

3) Hacerla comprensible. Frecuentemente una respuesta parece perfectamente comprensible al encuestador u observador, pero no lo es para otro. El contexto en el cual la conducta o respuesta ocurre es conocida por el encuestador, pero no por el codificador, quien de este modo no puede ver exactamente lo que el sujeto hizo, o comprender lo que su respuesta significa. Preguntas sistemáticas al encuestador u observador disipan confusiones y ambigüedades, mejorando considerablemente la calidad de la codificación.

4) Hacerla consistente. Marcadas inconsistencias en una cédula de entrevista u observación dada no sólo crean problemas en la codificación: pueden indicar errores en la recolección y registro de los datos... puede ser deseable tomar contacto nuevamente con el entrevistado si el punto es importante para el análisis de los datos.

5) Uniformarla. En conjunto, de adecuadas instrucciones a los encuestadores u observadores resultarán procedimientos uniformes en la recolección y registro de los datos; sin embargo, es necesario verificar la uniformidad con que estas instrucciones fueron seguidas.

Tiempo necesario para la codificación. Errores residuales

Antes de terminar este trabajo parece de interés transcribir algunos párrafos de dos puntos que Chevy estudia: a) Tiempo necesario de la codificación; y b) Errores residuales.

Tiempo necesario de la codificación

Dice Chevy:,¹⁶

Es pues absolutamente necesario, cuando se establece el plan de una encuesta importante, que se pueda evaluar bien, antes de la encuesta, y con una precisión suficiente, el tiempo que será necesario para asegurar la codificación de los cuestionarios, el costo de cuya operación representa una parte importante del presupuesto global de una encuesta. Si esta evaluación previa ha sido demasiado optimista, podemos encontrarnos con muy poco dinero cuando se trate de acabar la elaboración. Es necesario organizar la codificación con el mayor cuidado para efectuarla tan rápida y económica como sea posible. Se utilizarán a este efecto todos los medios que permitan mejorar el rendimiento de las codificadoras: división del trabajo, circulación de los documentos, equilibrio de los diferentes grupos previstos para evitar superposiciones, consignas claras y precisas, reenvío de casos difíciles o dudosos a los cuadros especialmente formados, etcétera.

Errores residuales

A este respecto Chevy señala:

A pesar de todo el cuidado puesto en las diferentes operaciones que hemos descrito, los cuestionarios codificados llevarán todavía un cierto número de errores de diversos orígenes: los errores que comportan las respuestas de las encuestas y que no hubieran sido eliminados por la verificación de los cuestionarios. Los errores de codificación que hubieran escapado a su verificación. Debe notarse a este propósito que un cierto número de estos errores no tienen ninguna posibilidad de ser corregidos si la verificación de la codificación ha sido hecha por sondeo. Estos errores sin embargo pueden ser clasificados en diferentes categorías según su frecuencia y según la influencia que tuvieron en los resultados de las clasificaciones.

a) Cierta número de esos errores pasará siempre inadvertido porque nada podrá revelar su existencia, ni permitir su localización. Será el caso de errores —a condición de que sean poco numerosos— que en los cuadros de resultados afectarán en más o en menos los casos donde deben figurar normalmente números relativamente importantes. Ningún control de verosimilitud permitirá detectarlos. Por lo demás, su presencia se traduce por errores relativos muy débiles, que tienen poca importancia.

b) No será igual cuando los errores, aun poco numerosos, afecten casos donde deben normalmente encontrarse números muy débiles. Un número absoluto de

errores muy poco importante puede entonces doblar, triplicar o, al contrario, reducir a cero al número que debería figurar en la casilla afectándolo en consecuencia con un error relativo considerable.

c) Ciertos errores que fueron cometidos muy frecuentemente porque tomaron un carácter sistemático podrán crear una anomalía muy visible en los cuadros estadísticos de resultados, lo cual proveerá, para ciertos casos, un número de unidades que aparecerá muy poco verosímil.

d) En fin, otros errores tendrán como resultado hacer figurar un pequeño número de unidades en una casilla de un cuadro donde no debería encontrarse ninguno, en razón de la naturaleza misma de las cosas, porque hay incompatibilidad fundamental entre los valores de los dos caracteres de los cuales esta casilla representa el cruzamiento.

Supongamos, en efecto, que una codificadora ha atribuido a una cierta profesión individual el número de código de otra profesión. Si este error es raro y no ha sido descubierto por la verificación, peligra no poder ser descubierto ulteriormente (caso a), a menos que tuviera por resultado hacer aparecer un pescador, por ejemplo, en una región donde esta profesión individual no tiene ninguna razón de estar representada (caso d).

Pero, si se trata de un error sistemático de codificación cometido numerosas veces, tendrá por efecto engrosar en gran medida el número de unidades que figuran en la casilla que corresponde a la profesión efectivamente codificada y reducir otro tanto el número de unidades que figuran en la casilla de la profesión verdadera (caso c). En el curso de un control de verosimilitud de los resultados, estos dos números podrán aparecer muy sospechosos, el primero por su importancia, el segundo por su debilidad.

Por otra parte, si la situación matrimonial 4 (divorciado) ha sido atribuida a una muchacha de 12 años, o aun si el *status* profesional 2 (empleador) ha sido dado a una persona cuya categoría socioprofesional ha sido codificada 60 (capataz), aparecerá una unidad en una casilla de los cuadros:

- repartición según la edad y el estado matrimonial, en el primer caso;
- repartición según la categoría socio-profesional y el *status*, en el segundo caso;

que debería permanecer absolutamente vacío porque hay incompatibilidad fundamental entre las modalidades de los caracteres en cuestión (caso d).

El tratamiento de otras cuestiones relativas a la codificación, tales como la construcción de índices y coeficientes, diversas modalidades que la codificación adopta en el caso de otros instrumentos de relevamiento distintos del cuestionario (análisis de contenido, entrevista focalizada, técnicas sociométricas, etc.), demandaría un esfuerzo de estudio que no pretendemos explorar aquí. Si bien resultaría interesante, puesto que son temas poco transitados en la literatura sociológica.

De cualquier manera, en el apéndice se discuten algunos puntos conectados con la

codificación y no mencionados en el presente trabajo.

Sin duda, los problemas tratados en dicho apéndice, muy sumariamente, admiten una elaboración más pormenorizada y un nivel de profundidad mayor. Pero lo que se pretendió con él es dejar abierta la posibilidad de ulteriores desarrollos sobre los temas a que se hace referencia.

APÉNDICE

1) La utilización de símbolos numéricos

Como ya se ha dicho, los símbolos numéricos no son de empleo obligatorio en las tareas de codificación; puede admitirse que los símbolos usados sean de naturaleza muy diversa. Además de los números, pueden emplearse letras, signos algebraicos, determinados cortes o perforaciones en una tarjeta, colores, rayas de diversos colores, etcétera.

Si debemos codificar sexo, por ejemplo, podemos adoptar como código el siguiente: en la tarjeta a la que trasladaremos los datos del instrumento de relevamiento haremos un corte en la esquina superior derecha para los varones y no haremos ninguno para las mujeres. Este simple artificio nos permite, por medio de símbolos muy sencillos, proceder a la codificación de la variable sexo, así como una rápida individualización, e incluso separación, de las tarjetas que caen en uno u otro ítem. Obsérvese que en los sanatorios las fichas de los recién nacidos son impresas en cartulinas de distinto color, según el sexo. Esto también implica que estamos clasificando al recién nacido en una de las dos clases en que se divide la variable sexo y de este modo la estamos codificando por medio de un determinado color; vale decir, estamos empleando símbolos cromáticos para codificar.

El elenco de símbolos puede ser aún mayor; pero los ejemplos mencionados, al parecer, alcanzan para mostrar algún aspecto nuevo referido a los símbolos en la codificación.

2) Codificación y cuadros bidimensionales

En este párrafo se pretende mostrar, en líneas muy generales, algún aspecto de la utilidad que la codificación puede prestar en la tarea de tabulación de datos para la construcción de cuadros bidimensionales.

Supongamos que en el plan de explotación de una investigación está previsto el cruzamiento de dos variables dicotomizadas: sexo y adhesión a uno de dos partidos políticos. En lugar de relevar los datos de las dos variables por separado, diseñamos las dos preguntas conjuntamente:

Sexo y adhesión a partido político

Masc. Partido A	☐	1
Masc. Partido B	☐	2
Fem. Partido A	☐	3
Fem. Partido B	☐	4
Masc. Sin respuesta	☐	5
Fem. Sin respuesta	☐	6

Cada uno de los números de código asignados a cada respuesta combinada corresponde a cada uno de los casilleros del cuadro a construir en la tabulación de los datos. Dicho cuadro sería como sigue:

	<i>M</i>	<i>F</i>
<i>A</i>	1	3
<i>B</i>	2	4
<i>S/R</i>	5	6

Al elaborarse los datos obtendríamos directamente los resultados correspondientes a dos preguntas, lo que facilitaría en buena medida los trabajos de tabulación.

Por cierto que el ejemplo es sumamente sencillo y pueden pensarse soluciones para el cruzamiento de más de dos variables y no solamente dicotomizadas. Pero parece suficiente el ejemplo si de lo que se trata es de presentar una posible utilización de la codificación más allá de su empleo corriente.

3) Codificación y escalas

Aunque la construcción de escalas para su utilización en las investigaciones sociológicas está en sus inicios, puede hacerse alguna referencia a la relación entre escalas y codificación. Esta relación es evidente cuando se trata de dimensiones claramente cuantitativas (o cuantificables), a saber, ingresos, edad, años que se desempeña una ocupación, etc. Si las respuestas a una pregunta sobre ingresos las categorizamos en ingresos de 0 a \$ 1 000, de \$ 1 000 a \$ 2 000 y de \$ 2 000 y más, y le asignamos al ítem primero el número 1, al que sigue el 2 y al último el 3, cada número de código corresponderá a un intervalo de la escala; la posición de un sujeto o un grupo de ellos estará dado por el número de código que le corresponda.

El problema es más complejo si queremos indagar acerca de la posible utilización de la técnica de la codificación referida a aquellas dimensiones más difíciles de cuantificar, por ejemplo, frente a la dimensión poder. La igualación de los intervalos de la escala, la determinación del cero son problemas que se plantean frente a este tipo de dimensiones y, por lo tanto, resulta más problemático usar la codificación provechosamente en estos casos. De cualquier modo esta posibilidad siempre existe, y su empleo, en el futuro, puede resultar en un considerable ahorro de esfuerzo en los trabajos de elaboración.

4) *Codificación y tabulación manual*

Transcribiremos lo que en *Manual de métodos para la elaboración de datos* se dice al respecto:

Codificación. Es útil describir algunas técnicas comunes de codificación empleadas en el sistema de separación manual de tarjetas perforadas. Puede usarse el fichero para seleccionar las tarjetas, para su arreglo en un cierto orden o para ambos propósitos. El tipo de código empleado está considerablemente influido por el uso principal que se dé al fichero. Los códigos deberán ser fáciles de usar y deben permitir el logro del máximo de rapidez y de eficiencia en las operaciones de separación. A continuación se describen algunos de los sistemas importantes de codificación:

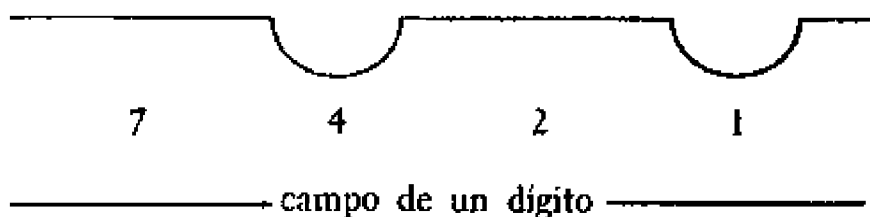
a) *Codificación directa o simple por ranuras.* Se asigna un orificio a cada clase requerida. Así, el número de clases que pueden ser ranuradas en una tarjeta será igual al número de orificios ranurables en dicha tarjeta (sin contar las posiciones de las esquinas, que normalmente no se ranuran). Por lo tanto, es limitado el número de características que pueden codificarse en una tarjeta por medio de este método. Sin embargo, cuando la selección es simple y la tarjeta tiene lugar suficiente para todas las clases requeridas, es preferible usar la codificación directa.

b) *Codificación combinada.* En este método de codificación generalmente se asigna más de un orificio, en una combinación adecuada, a cada clasificación, dígito o letra. Usando códigos combinados y también codificando adecuadamente las clasificaciones, puede aumentarse considerablemente la cantidad de grupos que es posible ubicar en una tarjeta. Generalmente se usa este método de codificación cuando las características codificadas exceden el número total de orificios ranurables. Con una codificación combinada, el fichero puede ordenarse en una secuencia numérica, con menos pasadas de la aguja separadora (es decir, manipulando las tarjetas un menor número de veces) que las necesarias con un código directo. Por esta razón, este método de codificaciones es útil cuando las tarjetas tienen que ser ordenadas en una secuencia numérica o alfabética.

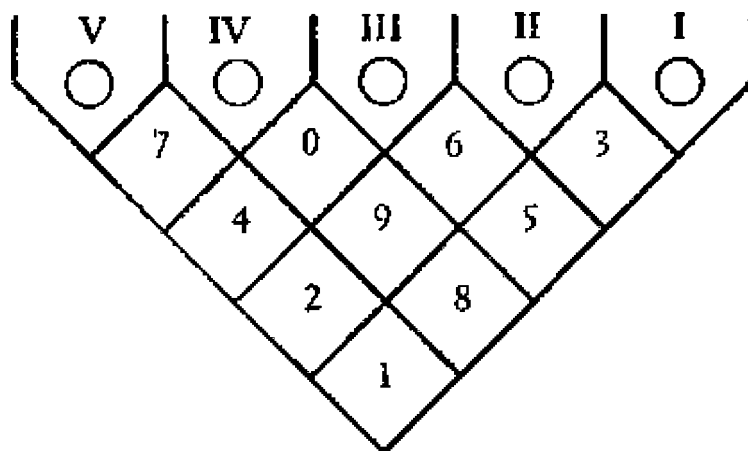
La combinación de los códigos, generalmente empleada para la codificación numérica y alfabética, puede derivarse de una interesante propiedad de la serie geométrica: 1, 2, 4, 8, 16, ..., con término general 2^n . Podrá notarse que las

combinaciones de diferentes términos de esta serie dan todos los enteros positivos; y cualquier término es mayor en una unidad que la suma de todos los otros términos anteriores. Los cuatro primeros (1, 2, 4, 8) dan, en combinaciones diferentes, todos los enteros positivos desde 1 hasta 15; por ejemplo, el número 7 se obtiene sumando 1, 2 y 4. Cualquier otro número puede obtenerse en una, y solo una, forma, es decir que hay solamente una forma de obtener 13, usando 8, 4 y 1.

En forma similar 1, 2, 4, 7 dan todos los enteros positivos desde 1 hasta 9, combinando *no más de dos* de cualquiera de estos números. Esto elimina el que 7 sea obtenido de 1, 2 y 4. Cualquier entero entre 1 y 9 puede situarse, por tanto, en cualquier grupo de 4 orificios consecutivos, asignándole a cada orificio uno de los cuatro valores y ranurando uno o dos de estos dos orificios. El cero queda representado no ranurando ninguno de ellos, o ranurando 4 y 7. De esta manera, en lugar de necesitarse 9 orificios para representar los enteros entre 1 y 9, son suficientes solamente 4. La siguiente figura ilustra un “cinco” ranurado en el campo de las unidades.



Si se desea codificar en tal forma que cada dígito pueda seleccionarse solamente con dos pasadas de la aguja, entonces es necesario asignar para cualquier dígito dos orificios para ranurar. Todo lo que hay que hacer para tener los dígitos entre 1 y 9 es sumar dos números del conjunto 0, 1, 2, 4, 7, mientras que el cero puede obtenerse ranurando 4 y 7. De simples reglas de combinación de álgebra, se sabe que de cinco orificios pueden obtenerse 10 combinaciones diferentes tomando dos orificios para cada combinación. Por tanto, cualquiera de los 10 dígitos (de 0 a 9) puede perforarse en cinco orificios, ranurando dos orificios, y puede emplearse un esquema triangular o, enrejado para mostrar el grupo de dos, de entre los cinco orificios, que debe ranurarse para obtener cada dígito particular. A continuación se muestra un esquema ilustrativo.



Los dos orificios que deben ranurarse para cualquier dígito son aquellos en los cuales terminan los dos conjuntos de celdas alineadas del enrejado, que contienen al dígito. Así, por ejemplo, los orificios a ranurar para 9 son II y IV, como se muestra por medio de las flechas. Si se emplean códigos visuales, en forma de enrejado, entonces la asignación de los dígitos a las diferentes celdas requiere algún cuidado, pues la clasificación del archivo según una secuencia numérica puede presentar algunos problemas. El modelo presentado arriba ha sido diseñado para facilitar la secuencia de la separación, haciendo corresponder I, II, III, IV y V, respectivamente, con 1, 2, 4, 7, 0.

5) Codificación y cuestionario

Veremos a continuación, a través de las palabras de Chevy,¹⁷ algún tipo especial de cuestionario en el que la codificación adopta modalidades diferentes de las vistas en el trabajo:

Conviene mencionar dos modalidades de presentación de cuestionarios que permiten hacer leer por una máquina las respuestas a las preguntas. En el primero, sistemas *mark sensing* (IBM) y magneto-lectura (*Bull*), la máquina que lee las respuestas las traduce en perforaciones en una tarjeta perforada normal que en seguida es explotada por las máquinas clásicas (clasificadora, tabuladora, etc.). En el segundo, llamado FOSDIC, de invención más reciente, las respuestas son traducidas mecánicamente, y sin la mediación de tarjetas perforadas, en signos sobre bandas magnéticas que serán explotados por calculadoras electrónicas.

Solamente transcribiremos la explicación de Chevy referente a la primera modalidad:

Los cuestionarios son impresos en tarjetas perforadas del formato clásico de 80 columnas (187.3 × 82.6 mm), o tarjetas de formato más grande. Las preguntas son dispuestas en columnas yuxtapuestas y son transcritas con la ayuda de un lápiz especial en los lugares previstos a este efecto en la columna reservada a cada

pregunta (de ahí el nombre de “grafitado” dado a esta transcripción). En el sistema *mark sensing* el grafitado consiste en marcar con un trazo el eje mayor de un óvalo. En el sistema *Bull* de magneto-lectura, los lugares donde debe hacerse el grafitado son pequeños cuadrados que llevan en su interior una cifra o una letra en trazo ligero y el grafitado consiste en reescribir con un lápiz especial esa cifra o letra, o, más simplemente, trazar las diagonales del cuadrado. En cada columna reservada a una pregunta, habrá tantos óvalos o rectángulos como respuestas posibles a esta pregunta. Si la respuesta debe expresarse por un número, habrá tantas columnas de óvalo o de rectángulos marcados de 0 a 9 como números de cifras significativas. Ejemplos: a) La pregunta “sexo” se presentará de la manera siguiente:

IBM	BULL
Sexo	Sexo
M	M
	F
F	

y el registro de la respuesta consistirá en marcar con un trazo el eje mayor del óvalo pertinente o trazar las diagonales del rectángulo *M* o *F*.

b) La pregunta “año de nacimiento”, en la inteligencia de que no se retendrá nada más que las dos últimas cifras, se presentará de este modo:

IBM		BULL	
Año de nacimiento		Año de nacimiento	
0	0	0	0
		1	1
1	1	.	.
		.	.
.	.	.	.
.	.	.	.
9	9	9	9

y la inscripción del año 1917 se hará marcando el óvalo (o el rectángulo) 1 en la columna de la izquierda y el óvalo (o el rectángulo) 7 en la columna derecha. Las tarjetas-cuestionarios de este modo grafitadas pasan en seguida por una máquina

reproductora, provista de un dispositivo especial que lee los trazos de lápiz y los traduce inmediatamente en perforaciones en la misma tarjeta.

VI. ESCALAS PARA LA MEDICIÓN DE ACTITUDES

JORGE PADUA
INGVAR AHMAN

EL PROPÓSITO de este capítulo es presentar una serie de escalas —conocidas como escalas para la medición de actitudes— pero que de hecho pueden ser utilizadas para mediciones de otras variables.

La idea principal es posibilitar al lector la construcción de cada una de ellas, indicando los diferentes pasos con los mayores detalles. Como en los capítulos anteriores, recomendamos al lector interesado en el tema recurrir a la bibliografía que detallamos al final del capítulo.

Previamente a la presentación de cada tipo de escala conviene realizar una serie de aclaraciones que son importantes tanto para la construcción de las escalas como para la interpretación de sus resultados:

1) La primera disquisición tiene que ver con el problema de la *medición*. Si bien en otro capítulo de este libro presentamos los diferentes niveles de medición en forma más extensa, interesa aquí dar un resumen sobre lo que se entiende por “medir” en ciencias sociales. La medición, de hecho, corresponde a una serie de teorías conocidas como *niveles de medición*: detrás de cada uno de los distintos niveles de medición están operando una serie de principios logicomatemáticos, que van a determinar o no el isomorfismo entre un concepto y el nivel de medición apropiado; es decir, el problema de la medición es expresado aquí como un problema en el cual se busca que el modelo matemático sea isomórfico con el concepto; esto es, que la “forma” del modelo sea idéntica a la “forma” del concepto. Si esto no ocurre (recuérdese que no hay isomorfismos “parciales”), estamos deformando el concepto (por ejemplo, si aplicamos niveles o modelos matemáticos intervalares a conceptos solamente operacionalizados a nivel ordinal es muy probable que los resultados ulteriores sean una consecuencia más del modelo matemático que del concepto en sí). Los niveles de medición más utilizados en ciencias sociales son:

a) Nivel nominal: la operación de medir consiste simplemente en la asignación de nombres o de números a distintas categorías. La función del número en este nivel de medición es muy elemental, ya que simplemente sirve para distinguir diferentes categorías. A este nivel se está “midiendo” cuando, por ejemplo, se hace la distinción entre varón y mujer o entre católico, protestante, judío, mahometano o de otra religión. La operación de medición consistiría entonces en referir la observación a una clase o categoría, para luego contar cuántas frecuencias caen dentro de cada categoría. Uno no

puede hacer legítimamente ninguna afirmación que vaya más allá de las diferentes distinciones.

b) Nivel ordinal: en este nivel de medición uno está en condiciones de distinguir entre diferentes categorías y de poder afirmar si una categoría posee en *mayor*, *menor* o *igual* grado el atributo que estamos midiendo. La escala de jerarquía militar en el ejército es un buen ejemplo: un sargento tiene *menos* autoridad y es *diferente* de un teniente; éste a su vez tiene *menor* grado de autoridad que un capitán. Uno puede ordenar entonces las categorías sargento, teniente, capitán, con respecto a autoridad de la siguiente manera:

Capitán > Teniente > Sargento

Incluso puede llegarse a establecer comparaciones con distribuciones de autoridad para distintas categorías en la marina o las fuerzas aéreas, etcétera.

c) Nivel intervalar: existe aún mayor precisión que en los anteriores, ya que no solamente podemos categorizar y establecer relaciones de mayor, menor o igual, sino además calcular la distancia entre los intervalos o categorías. Obsérvese que en el caso de las mediciones de nivel ordinal uno puede afirmar que un sargento tiene menos autoridad que un teniente, y que un capitán tiene más autoridad que éste; sin embargo, no estamos en condiciones de precisar *cuánto más* o *cuánto menos*. Mediante la adjudicación de un cero arbitrario en esta escala podemos especificar la distancia entre esas categorías.

d) Nivel racional: es la forma de medición que utiliza valores cero absolutos, y que nos permiten establecer diferencias entre cualquier par de objetos a un máximo de precisión. A esta escala pertenecen el sistema métrico y el de pesos, por ejemplo. La diferencia entre el nivel intervalar y el nivel racional es que, por ejemplo, en la medición de la temperatura en que se utiliza escala intervalar, no es posible afirmar correctamente que cuando se registran 40 grados de temperatura estamos sintiendo el *doble* de calor que cuando teníamos 20 grados; por el contrario, si nos desplazamos 40 kilómetros en línea recta sí podemos afirmar que hemos duplicado la distancia recorrida al desplazarnos 20 kilómetros.

La mayor parte de las escalas de medición de actitudes que vamos a describir se encuentran comprendidas entre la escala ordinal y la escala intervalar.

2) La segunda disquisición a considerar también ha sido tratada en otro capítulo. Tiene que ver con la dimensionalización de los conceptos. (Ver “El proceso de investigación”.) Recordemos solamente que muchas de las variables con las que trabajan los científicos sociales son complejas y están compuestas de una serie de dimensiones o atributos. Por ejemplo, la variable “religiosidad” puede ser concebida como compuesta de tres dimensiones: dogmatismo, misticismo, ritualismo, cada una de las cuales tiene distintos indicadores. O la del *status* socioeconómico, que tradicionalmente se dimensionaliza en ocupación, educación e ingreso. El investigador en general espera que las dimensiones estén intercorrelacionadas, es decir, uno espera que una persona con alto grado de religiosidad manifieste valores altos en las dimensiones ritualismo, misticismo y

dogmatismo. Idénticamente en el caso del *status* socioeconómico es de esperar que la educación se correlacione estrechamente con la ocupación y el ingreso. Sin embargo, existen casos en los cuales se producen “inconsistencias”: sujetos con alta educación y bajo ingreso; alta ocupación y baja educación, etc.; o en el caso de religiosidad, sujetos que responden positivamente a ítems de dogmatismo y negativamente a ítems de ritualismo. En el caso de escalas es posible construir *escalas multidimensionales* o *escalas unidimensionales*.

3) La tercera digresión es parte ya del discurso propio de la construcción de escalas para la medición de actitudes y tiene que ver con los modos en que se incluyen o eliminan los ítems de una escala. Una vez que los ítems o los juicios de actitud han sido formulados o recolectados, los métodos utilizados para su inclusión en las escalas son:

a) Uso de jueces que no responden a los ítems en términos del grado de acuerdo o desacuerdo que se tenga con ellos, sino en el grado de validez que el juez otorgue al ítem o juicio en relación a la variable. Es decir que los jueces son utilizados aquí para determinar el valor que el investigador va a asignar al ítem sobre un continuo psicológico. Una vez que los juicios o ítems tienen asignado un valor, se aplican a los sujetos para que ellos *sí* expresen su grado de acuerdo o desacuerdo. Los puntajes definitivos serán computados a partir del valor dado por los jueces. Las escalas que desarrollaremos más adelante y que utilizan este sistema son: la del método de intervalos aparentemente iguales de Thurstone, la de intervalos sucesivos y la de comparación por pares.

b) El método de las respuestas directas con los ítems o juicios. Este método no requiere el conocimiento previo de los valores de escala, sino que los puntajes se determinan en función de las respuestas. No es necesaria la utilización de jueces en el sentido expresado en el párrafo anterior. El método a examinar aquí será el del análisis de escalograma y el de la escala Lickert.

c) Finalmente, en la combinación de método de respuesta y de uso de jueces, el método a examinar será el de la técnica de la escala discriminatoria (*scale-discrimination technique*).

4) La cuarta digresión tiene que ver con los problemas de confiabilidad y validez, y con que los resultados en el test pueden sufrir variaciones de tres tipos:

a) Variaciones en el instrumento: los errores producto de instrumentos “mal calibrados” (problema de validez).

b) Variaciones en los sujetos: dicho en otros términos, en dos aplicaciones distintas en el tiempo, el sujeto produce resultados distintos (problemas de confiabilidad).

c) Variaciones simultáneas en los sujetos y en el instrumento: por instrumento “válido” entendemos aquel que mide efectivamente lo que se propone medir; mientras que por “confiable” entendemos que mide siempre de la misma manera.

5) La quinta digresión se relaciona estrechamente con la tercera y la cuarta, y se refiere a las formas en que sean clasificadas las escalas según estén centradas en los sujetos, en el

instrumento o en ambos:

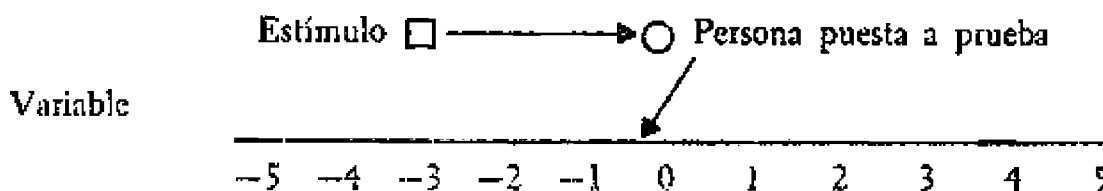
a) El enfoque centrado en el instrumento, llamado *stimulus-centered approach* (Torgerson), es aquel en el cual la variación sistemática en la reacción de los sujetos al estímulo es atribuida a diferencias en éste.

b) En el enfoque centrado en el sujeto (*subject-centered approach*) la variación al estímulo es atribuida a diferencias individuales en los sujetos.

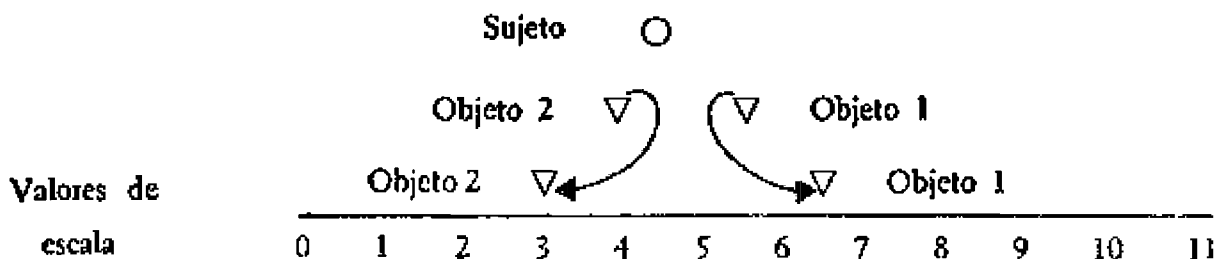
c) En el enfoque centrado en la respuesta (*response approach*) la variabilidad de reacciones al estímulo se atribuye tanto a las reacciones de los sujetos como al instrumento.

6) Otra forma de clasificar los instrumentos o escalas es la propuesta por Coombs,¹ de tests de Tipo A y tests de Tipo B:

a) El test de tipo A: sirve para determinar algunas propiedades de un sujeto o de un objeto en el medio de una persona. A través de su conducta en la situación de prueba (es decir, su rendimiento en un ítem específico de la escala) el sujeto, consciente o inconscientemente, se sitúa en una posición a lo largo del continuo en la variable que la escala está midiendo. Gráficamente:



b) El test de tipo B: se usa para determinar las propiedades de un objeto. El sujeto indicaría entonces la posición de uno o más objetos (ítems en nuestro caso) a lo largo de la variable. En la construcción de escalas, el test de tipo B se utiliza por ejemplo cuando se pide a jueces que ubiquen el valor que un determinado ítem tiene en cierta variable (ver escala Thurstone). La situación gráfica es la siguiente:



Ejemplo: Si usted le pregunta a un número de personas cuál es el partido político que tiene el programa más conservador en Argentina, el ítem o pregunta puede ubicarse como un test de tipo B. Si la pregunta fuera: ¿Cuál es el partido con el que usted más simpatiza?, se trataría de un test de tipo A.

7) La última digresión tiene que ver con los ítems. Por ahora podemos definir el ítem

como una frase o juicio indicador de la variable que estamos tratando de medir. La formulación del ítem en el caso en que estemos trabajando con cuestionarios es siempre un juicio, nunca una pregunta o una interrogación. Para relacionar los ítems con los niveles de medición mencionados en la primera digresión tomemos los siguientes ejemplos:

A) ¿Entre cuáles de los siguientes grupos de altura se ubica usted? (marque con una cruz la que corresponda)

☐ 120-139 cm

☐ 140-159 cm

☐ 160-179 cm

☐ 180-199 cm

☐ 200-219 cm

B) Si usted tuviera que definir sus intereses en los asuntos políticos, ¿diría que está muy interesado, algo interesado, ni desinteresado ni interesado, desinteresado o muy desinteresado?

☐ Muy interesado

☐ Interesado

☐ Indiferente

☐ Desinteresado

☐ Muy desinteresado

C) De los siguientes, ¿cuál es el diario que lee con mayor frecuencia?

- ☐ *Excelsior*
- ☐ *Novedades*
- ☐ *Últimas Noticias*
- ☐ *El Herald*
- ☐ *El Día*
- ☐ *El Sol de México*
- ☐ *Ovaciones*
- ☐ *La Prensa*
- ☐ Otro: (especificar)

El ejemplo A señala una variable medida a nivel racional;
 el ejemplo B corresponde a un nivel de medición ordinal y
 el ejemplo C corresponde a un nivel de medición nominal.

Ahora bien, si tratamos de medir la altura de una persona, como en el caso del ejemplo A, es evidente que podemos hacer la medición de varias maneras. Por un lado podríamos tener diferentes varas con diferentes alturas: una vara de 175, otra de 176, otra de 177, etc. Supongamos que nos llega un sujeto de 178 cms de altura; le aplicamos la vara de 160; la rechazamos diciendo que el sujeto es más alto, y seguimos aplicando varas hasta que llegamos a la correcta. En este caso estamos tratando con lo que se denomina ítems *acumulativos*, que en el cuestionario autoadministrado aparecerían de la siguiente forma: (las marcas aparecen para el caso de nuestro sujeto que mide 178 cm)

<i>ítem</i> 1: ¿Tiene Ud. más de 170 cm de altura?	Sí ☒	No ☐
<i>ítem</i> 2: ¿Tiene Ud. más de 171 cm de altura?	Sí ☒	No ☐
<i>ítem</i> 3: ¿Tiene Ud. más de 172 cm de altura?	Sí ☒	No ☐
<i>ítem</i> 4: ¿Tiene Ud. más de 173 cm de altura?	Sí ☒	No ☐
<i>ítem</i> 5: ¿Tiene Ud. más de 174 cm de altura?	Sí ☒	No ☐
<i>ítem</i> 6: ¿Tiene Ud. más de 175 cm de altura?	Sí ☒	No ☐
<i>ítem</i> 7: ¿Tiene Ud. más de 176 cm de altura?	Sí ☒	No ☐
<i>ítem</i> 8: ¿Tiene Ud. más de 177 cm de altura?	Sí ☒	No ☐
<i>ítem</i> 9: ¿Tiene Ud. más de 178 cm de altura?	Sí ☐	No ☒
<i>ítem</i> 10: ¿Tiene Ud. más de 179 cm de altura?	Sí ☐	No ☒
<i>ítem</i> 11: ¿Tiene Ud. más de 180 cm de altura?	Sí ☐	No ☒
Etcétera.		Etc.

Si una persona ha contestado a una de las preguntas con un “Sí”, sabemos que todas las preguntas que están por “arriba” tienen también, por derivación, una respuesta “Sí”. De la misma manera, si una persona ha respondido “No”, todas las preguntas que estén por “debajo” tendrán una respuesta similar.

Una alternativa diferente de presentar la pregunta es a través de ítems *diferenciados*,² en la que, ante una serie de ítems, la persona entrevistada marcará solo una o algunas alternativas. En un cuestionario, la pregunta sería realizada de la siguiente manera:

<i>ítem</i> 1: ¿Está Ud. entre los 170-174 cm de altura?	Sí	☒
<i>ítem</i> 2: ¿Está Ud. entre los 175-179 cm de altura?	☒	No
<i>ítem</i> 3: ¿Está Ud. entre los 180-184 cm de altura?	Sí	☒
<i>ítem</i> 4: ¿Está Ud. entre los 185-190 cm de altura?	Sí	☒

Un ítem incluido en una serie de ítems diferenciados debe ser formulado de manera tal que sea respondido con un “Sí” únicamente por las personas que tienen una posición fija a lo largo de la variable. Un ítem incluido en una serie de ítems acumulativos debe ser formulado de manera tal que solamente pueda ser respondido con un “Sí” a un lado de una determinada posición a lo largo de la variable investigada.

Trabajando con la altura de las personas, la diferencia entre estos dos tipos de serie de

ítems aparece como trivial. Sin embargo, en el tratamiento de las escalas para medición de actitudes, su diferenciación es importante. Consecuentemente, en cada tipo de las escalas a examinar (Lickert, Thurstone, Guttman, comparación por pares) tomaremos en cuenta la especificación de esta perspectiva.

ESCALAS PARA LA MEDICIÓN DE ACTITUDES

Vamos a utilizar la definición de escalas hecha por Stoufer:³

Se dice que existe escala cuando a partir de una distribución de frecuencias multivariada de un universo de atributos es posible derivar una variable cuantitativa con la cual caracterizar los objetos de un modo tal que cada atributo sea una función simple de aquella variable cuantitativa.

Una escala es una forma particular de *índice*, aunque aquí utilizaremos una serie de procedimientos objetivos para la selección de ítems destinados a controlar los errores producto de la subjetividad del investigador. (Ver el capítulo sobre “Conceptos, indicadores e índices”.) Construir una escala implica una serie de procedimientos mediante los cuales —de acuerdo a distintas reglas— se seleccionan ítems y se adjudican números a un conjunto de ítems (juicios o sentencias), número que va a expresar la intensidad que un sujeto o grupo de sujetos manifiestan en la variable.

Las actitudes en el contexto individual representan un estado mental que es un puente entre estados psicológicos y objetos exteriores. Kretch y Cruschfield ⁴ sostienen a este respecto que se puede definir una actitud “como una organización durable de procesos motivacionales, emocionales, perceptuales y cognitivos con respecto a algún aspecto del mundo del individuo”. Las actitudes serían entonces procesos claves para entender las tendencias del individuo en relación con objetos y valores del mundo externo, aunque esas tendencias no son estáticas y, como las de Newcomb,⁵ las actitudes representan un residuo de la experiencia anterior del sujeto. Las actitudes perdurarían en el sentido en que tales residuos “son trasladados a nuevas situaciones, pero cambian en la medida en que nuevos residuos son adquiridos a través de experiencias en situaciones nuevas”.

Las actitudes serían entonces tendencias a actuar con respecto a alguna entidad especificable (Newcomb); o como quieren Thomas y Znaniecki: ⁶ “la tendencia individual a reaccionar, positiva o negativamente, a un valor social dado”.

Las actitudes medidas por escalas deben interpretarse en términos analíticos no como “hechos”, sino como “síntomas”. Existe una serie de conceptos relacionados con las actitudes; entre ellos detallamos los siguientes:

Creencia: actitudes que incorporan una cantidad importante de estructuración cognitiva. Las actitudes son *hacia* algo, mientras que las creencias son *en* o *sobre* algo.

Sesgo (bias): son actitudes o prejuicios débiles, basados en premisas incompletas, deducidas falsamente o preconcebidas. Por lo tanto son poco precisas y relativamente fáciles de cambiar.

Doctrina: son objetos estímulos elaborados, hacia los cuales el individuo manifiesta una actitud. Una doctrina (republicana, liberal, católica, comunista, etc.) describe específicamente las razones para adherencia; por lo tanto una doctrina se aprende.

Fe: implica una actitud con alta carga emocional o afectiva. El sistema de actitudes referentes a la fe describe una creencia fundamental y específica de la persona. La fe se ubica entre la creencia y la ideología.

Ideología: es un sistema cognitivo elaborado, que sirve para justificación de formas específicas de comportamiento, o como medio de racionalización. La ideología es concebida como un sistema lógico falso. La ideología se acepta como una fe.

Valor: en un sentido psicológico amplio, valores son marcos de referencia que sirven de guía o mapa para la evaluación de la experiencia y la conducta. Sistema de valores sería la organización elaborada y articulada de actitudes, que toman valencias positivas, negativas o neutras en cuanto a objetos, estímulos del ambiente y la relación de éstos con las metas de vida.

Opinión: son evaluaciones tentativas, no fijas, sujetas a cambio o inversión. Es decir, son menos fijas y no comprometedoras para el individuo.

ACTITUDES E INTERESES: ALGO ACERCA DE SU DIRECCIÓN E INTENSIDAD

Señalamos brevemente aquí algunas propiedades de las actitudes que tienen que ver con los problemas asociados a la medición de las mismas a través de escalas.

1) Dirección. La actitud puede ser positiva o negativa. En algunos casos es explicable hablar de una actitud neutral cuando el sujeto no es ni positivo ni negativo hacia el objeto.

2) Intensidad. La intensidad de la actitud es alta si el sujeto está fuertemente convencido de que la actitud es justificada y baja si el sujeto no piensa así. Al respecto, una medida aceptable es la frecuencia con la cual el sujeto usa fuertes expresiones que señalan un engrane emocional y en la medida en que él está preparado a argumentar en favor de su posición.

3) Estabilidad. La actitud es estable si permanece invariable por un periodo muy largo.

4) Fortaleza. Una actitud es fuerte si, por ejemplo, es difícil de cambiar a través de persuasión o propaganda, y suave si cambia fácilmente.

5) Importancia. La importancia de la actitud es grande si influye sobre la conducta de una persona en muchos campos de actividades.

6) Visibilidad, observabilidad o relevancia externa. Ésta es alta si es fácil concluir a partir de observaciones sobre la conducta del sujeto (manera de hablar, acciones) que él tiene cierta actitud.

7) Relevancia interna. Es detectable si en la experiencia del sujeto la actitud por sí misma tiene una parte importante en su mundo de referencia. La actitud de una persona hacia la religión puede tener baja relevancia externa y, al mismo tiempo, una alta relevancia interna.

8) Involucramiento del ego. Cuando una actitud llega a constituir una parte importante de la personalidad, o es algo adherido a ella.

9) Integración y aislamiento. De un sistema mayor de actitudes interrelacionadas (perspectiva de vida, ideologías). Las actitudes políticas o religiosas están frecuentemente incluidas en tal sistema complejo, mientras que una actitud hacia una cierta pasta de dientes es experimentada como aislada, sin relación lógica o psicológica con el complejo sistema total.

10) Especificidad o precisión. La imaginación es dirigida hacia la actitud. Ejemplo: un profesor de teología tiene una actitud más específica hacia la Iglesia católica que el "hombre de la calle".

11) Verificabilidad. La imaginación es igualmente dirigida hacia la actitud. El conocimiento es verificable, la fe, no; las opiniones ocupan un lugar intermedio y pueden a veces ser verificadas.

Debe notarse que muchas de las propiedades de las actitudes están correlacionadas. Alta intensidad frecuentemente significa estabilidad, fortaleza, alta relevancia interna y externa y alto grado de involucramiento del ego. Es importante observar los casos en los cuales las propiedades no están combinadas. En dichos casos, las distinciones hechas anteriormente son obvias.

Las pruebas ordinarias de actitud sólo admiten la medición de algunas de las propiedades mencionadas, principalmente dirección e intensidad. Para determinar las otras propiedades de la actitud podría ser útil una entrevista profunda del tipo entrevista focalizada.

Dirección e intensidad frecuentemente son vistas como una sola propiedad de la actitud. Este enfoque es justificado por lo siguiente: suponiendo que dirección e intensidad son dos variables independientes, si de cada una de ellas pueden obtenerse tres valores diferentes, las combinaciones siguientes son posibles:

CUADRO 1

<i>Intensidad</i> \ <i>Dirección</i>	<i>Positivo</i>	<i>Neutral</i>	<i>Negativo</i>
Fuerte	+++	000	---
Ni fuerte ni débil	+	00	--
Débil		0	

Una serie de investigaciones (entre otras la realizada por Edward H. Suchman) han mostrado, sin embargo, que una actitud neutral demuestra una baja intensidad. En la mayoría de los casos, las dos áreas rayadas podrían ser consideradas como inexistentes.

Los resultados son siete combinaciones, las cuales pueden ser ordenadas de la siguiente manera:

CUADRO 2. *Similitudes y diferencias entre algunos tipos de escalas para la medición de actitudes y en los enfoques Centrado en el sujeto, Centrado en el estímulo y Response-approach*

Similitudes	DIFERENCIAS		
	Centrado en el sujeto	Centrado en el estímulo	Response-approach
Los tres tipos de enfoques se usan en la construcción de escalas que permiten medir distinciones de grado, más que de cualidad entre individuos	Consiste en preguntarle a un sujeto su opinión acerca de un objeto para que él se ubique en determinado punto de la escala	Consiste en preguntarle dónde ubicaría un estímulo sobre un continuo que representa el objeto de estudio	El propósito puede ser escalar los sujetos, los estímulos, o ambos en relación a un objeto
	Se dan valores a los sujetos	Se dan valores a los estímulos	Se dan valores tanto a los sujetos como a los estímulos
	Se subrayan las diferencias individuales	Se tratan de eliminar las diferencias individuales	Se tienen en cuenta las diferencias en los sujetos y en los individuos
	Análisis de la varianza unidireccional	Análisis de la varianza unidireccional	Análisis de la varianza en dos direcciones
<i>Similitudes en las escalas Lickert-Thurstone-Guttman</i>	<i>Ejemplo: escala Lickert</i>	<i>Ejemplo: escala Thurstone</i>	<i>Ejemplo: escala Guttman</i>
Con pequeñas modificaciones los mismos ítems pueden utilizarse en los diferentes tipos de escala, ya que la diferencia entre las escalas no reside en la elección de los ítems, sino en la relación lógica existente entre ellos	Escala ordinal	Pretende ser intervalar	Escala ordinaria
	Escala aditiva	Escala diferencial	Escala acumulativa
	Pretende ser unidimensional, pero de hecho se mezclan dimensiones	<i>Idem.</i> a la Lickert	Es unidimensional, ésta es determinada por el coeficiente de reproductividad
Las tres emplean el método de jueces en la construcción de la escala	Los ítems se elaboran unidireccionalmente, procurando un 50% de ítems positivos y un 50% de ítems negativos	Los ítems se formulan de manera que $\frac{1}{2}$ sean positivos; $\frac{1}{2}$ negativos y $\frac{1}{2}$ neutros	Como es una técnica de análisis, los ítems se formulan de acuerdo a la técnica que se decida utilizar, Lickert o Thurstone
Estrictamente todas las escalas miden a nivel ordinal	Los ítems son acumulativos, por ello pueden estar aproximadamente en la misma posición de la escala	Los ítems son diferenciadores	Cada ítem tiene el carácter de acumulativo
En principio las tres son unidimensionales	La cantidad de ítems en la versión de los jueces varía de 30 a 50. En la versión final quedan entre 15 y 25	El número de ítems iniciales puede llegar a los 200. En la escala final pueden quedar un máximo de 50	La cantidad de ítem depende de la técnica de construcción
En el momento de aplicar la escala ya validada (versión final), aunque se hayan utilizado distintos enfoques en la construcción, todos los procedimientos van a estar centrados en el sujeto. Lo que subyace en todos los procedimientos es obtener conocimientos de las actitudes que forman parte de un contexto sociocultural	La selección de los ítems se realiza en base a su poder discriminativo	La selección de los ítems se hace en base al recorrido intercuartílico	La selección de los ítems se basa en su escalabilidad, que se determina en base al número de errores
	Las alternativas de respuesta en cada ítem pueden variar de 3, 4, 5 o más alternativas	El puntaje de los ítems se determina por el valor de la mediana de los puntajes asignados por jueces, sobre una escala que va de 1 a 11	El puntaje de los ítems varía de acuerdo a la técnica empleada
	No existe gradación de los ítems a lo largo de un continuo	Se presentan mezclados ítems positivos, negativos y neutros	Los ítems se ordenan en forma decreciente de acuerdo a su grado de dificultad

CUADRO 3

Semejanzas	DIFERENCIAS		
	Centrado en el sujeto	Centrado en el estímulo	Response-approach
	El número de jueces varía de 50 a 100. Estos deben tener características similares a las de los sujetos en el universo a estudiar	El número de jueces es entre 50 y 200. Se exige de ellos objetividad e información	El número de jueces depende de la técnica empleada en la construcción
	La consistencia interna de la escala se establece mediante el método de correlación por mitades (<i>split-half</i>)	La consistencia interna se basa en el cálculo de la correlación entre cada ítem con el puntaje total del test (<i>item-test</i>)	La consistencia interna está determinada por la escalabilidad, calculada en base a la diferencia CR-MMR
	A cantidad igual de ítems es más confiable que la escala Thurstone. La confiabilidad aumenta con el incremento en las alternativas de respuesta	Cuando se alcanza un número de 50 ítems es más confiable que la Lickert	Es más confiable debido a su unidimensionalidad
	Es más fácil y rápida de construir	Difícil de construir. Gasto fuerte en términos de tiempo y trabajo	Depende de la técnica utilizada en la construcción
	Escala correspondiente a test de Tipo A	Corresponde a un test de Tipo B	Corresponde a ambos tipos de test (A y B)
	En la escala final se presentan los ítems con la cantidad de alternativa idéntica a la de la versión de los jueces	En la escala final se presentan solamente dos alternativas de respuesta: acuerdo-desacuerdo	En la escala final se presentan los ítems en orden de dificultad creciente
	El puntaje máximo es igual al número de ítems multiplicado por el puntaje mayor en cada alternativa de respuesta; el puntaje mínimo es igual al número de ítems multiplicado por el puntaje menor en las alternativas de respuesta. Para la ubicación de individuos se pueden utilizar también en valores promedio	El puntaje máximo y mínimo en la escala se determinan por la sumatoria de ítems ponderados	Los sujetos se ubican en la escala en forma decreciente: desde los que contestaron en forma positiva a todas las preguntas hasta los que contestaron en forma negativa a todas las preguntas
	Con frecuencia la puntuación total de un individuo puede tener un significado poco claro, cuando se compara con otros individuos, ya que combinaciones distintas pueden producir el mismo resultado	<i>Idem.</i> a la Lickert	Los individuos que tienen un puntaje de actitud más favorable, deben también tener una actitud más favorable en cada ítem escalado
	El problema principal de la escala es la validez. Determinar cuándo una misma puntuación alcanzada por combinación de distintas categorías de respuesta tiene consecuencias para la interpretación del atributo en cuestión y cuándo no	El problema principal es el centro de la escala. Este no nos dice mucho acerca del significado que tiene el hecho de que un individuo ocupe una posición en el centro de la escala	En el caso de actitudes complejas, no es muy eficaz

El método del *summated ratings* de Lickert resulta de la suma algebraica de las respuestas de individuos a ítems seleccionados previamente como válidos y confiables. Si bien la escala es aditiva, no se trata de encontrar ítems que se distribuyan

uniformemente sobre un continuo “favorable-desfavorable”, sino que el método de selección y construcción de la escala apunta a la utilización de ítems que son definitivamente favorables o desfavorables con relación al objeto de estudio. El puntaje final del sujeto es interpretado como su posición en una escala de actitudes que expresa un continuo con respecto al objeto de estudio.

La escala Lickert es pues un *test de Tipo A*, ya que el sujeto, a través de su conducta en la situación de prueba, consciente o inconscientemente, se sitúa a lo largo de la variable. La escala Lickert es también una escala del tipo *centrada en el sujeto* (*subject-centered*): el supuesto subyacente es que la variación en las respuestas será debida a diferencias individuales en los sujetos. Veremos además, cuando examinemos los pasos en la construcción de la escala, que la escala inicial se administra a una muestra de sujetos que actuarán como *jueces*. (Esa muestra de sujetos debe ser representativa de la población a la que se aplicará la escala final.)

Finalmente, los ítems son seleccionados en base a su *poder discriminatorio* entre grupos con valores altos y con valores bajos en la variable. Es decir que lo que interesa es la coherencia, entendida en función de las respuestas.

A) La construcción de una escala Lickert

La construcción de una escala de este tipo implica los siguientes pasos: 1) Es necesario construir una serie de ítems relevantes a la actitud que se quiere medir. 2) Los ítems deben ser administrados a una muestra de sujetos que van a actuar como *jueces*. 3) Se asignan puntajes a los ítems según la dirección positiva o negativa del ítem. 4) Se asignan los puntajes totales a los sujetos de acuerdo al tipo de respuesta en cada ítem, la suma es algebraica. 5) Se efectúa un análisis de ítems. 6) Se construye con base en los ítems seleccionados la escala final. Examinemos en detalle cada uno de los pasos:

1) La construcción de los ítems

Los ítems que van a aplicarse a la muestra inicial de jueces, cuyo número debe ser entre 30 y 50. Para la construcción de los ítems deben tomarse en cuenta los siguientes criterios, que aparecen en Edwards: ⁷

- a) Evite los ítems que apuntan al pasado en lugar del presente.
- b) Evite los ítems que dan demasiada información sobre hechos, o aquellos que pueden ser interpretados como tales.
- c) Evite los ítems ambiguos.
- d) Evite los ítems irrelevantes con respecto a la actitud que quiere medir.
- e) Los ítems en la escala deben formularse según expresen actitudes o juicios favorables o desfavorables con respecto a la actitud. *No se trata* de elegir ítems que expresen distintos puntos en el continuo.
- f) Evite los ítems con los cuales todos o prácticamente nadie concuerda.
- g) Los ítems deben ser formulados en lenguaje simple, claro y directo.

h) Solamente en casos excepcionales exceda de las 20 palabras cuando formule el ítem.

i) Un ítem debe contener sólo una frase lógica.

j) Los ítems que incluyan palabras como “todos” “siempre”, “nadie”, etc. deben omitirse.

k) De ser posible, los ítems deben ser formulados con frases simples, y no compuestas.

l) Use palabras que el entrevistado pueda comprender.

m) Evite las negaciones, particularmente las dobles negaciones.

n) Combine los ítems formulados positiva y negativamente en una proporción aproximada a 50% — 50%.

Un sistema que puede ser aplicado para eliminar muchos ítems dudosos o que dan demasiados hechos es el siguiente: cada miembro del grupo de investigación responde a los ítems asumiendo primero una actitud positiva hacia la variable y luego responde como si tuviese una actitud negativa. Si la respuesta en ambos casos se ubica en la misma categoría, el ítem no es apropiado para incluir en la versión de los jueces.

Cada ítem es entonces un juicio o una sentencia a la cual el juez debe expresar *su* grado de acuerdo o desacuerdo. La graduación de acuerdos o desacuerdos varía en cantidad de alternativas que se le ofrece al sujeto; éstas pueden ser 3, 4, 5, 6 o 7 alternativas. En general, la decisión sobre la cantidad de alternativas a ofrecer dependerá no tanto de las “preferencias personales” del investigador, sino del tipo de investigación, del tipo de pregunta, del tipo de distribución de la variable, etc. Ejemplificaremos a continuación algunos ítems con sus respectivas alternativas.

Las siguientes afirmaciones son opiniones con respecto a las cuales algunas personas están de acuerdo y otras en desacuerdo. Indique, por favor (marcando con una X en el paréntesis correspondiente), la alternativa *que más se asemeja a su opinión*.

- | | |
|---|---|
| 1) Las mujeres no deberían meterse en política. | ☐ Muy de acuerdo
☐ De acuerdo
☐ Ni acuerdo ni desacuerdo
☐ En desacuerdo
☐ Muy en desacuerdo |
| 2) Leyendo lo que se publica en los diarios y las informaciones de radio y TV, es posible tener una idea acertada de lo que ocurre en la situación política mexicana. | ☐ Muy de acuerdo
☐ De acuerdo
☐ En desacuerdo
☐ Muy en desacuerdo
☐ Ni acuerdo ni desacuerdo |
| 3) Los manifiestos, proclamas y solicitudes que publican en los diarios los partidos políticos no informan sobre sus verdaderos propósitos. | ☐ Totalmente de acuerdo
☐ Medianamente de acuerdo
☐ Escasamente de acuerdo
☐ Escasamente en desacuerdo
☐ Medianamente en desacuerdo
☐ Totalmente en desacuerdo |
| 4) Vivo intensamente el presente sin pensar en el futuro. | ☐ Verdadero
☐ Falso |

Los ejemplos expresan ítems positivos y negativos; manifiestos y latentes y con distintas alternativas de respuesta. Los ítems son contruidos a partir de juicios que expresan alguna relación postulada a nivel de la teoría sustantiva, y de observaciones empíricas de afirmaciones de grupos o sujetos que pertenecen a grupos o asociaciones que manifiestan la propiedad que se quiere medir. Los ítems así pueden ser extraídos de libros, publicaciones y artículos que tratan teóricamente sobre el objeto que se quiere medir. También puede el investigador acudir a análisis de contenidos sobre discursos o manifiestos de individuos y asociaciones (por ejemplo, si se trata de medir radicalismo-conservadurismo, una fuente muy rica para la formulación de ítems son las declaraciones de grupos de interés: empresarios, grupos políticos de izquierda y de derecha, etc.). Otras estrategias para la construcción de ítems aparecen señaladas en la sección correspondiente a la escala Thurstone.

2) La administración de los ítems a una muestra de jueces

Una vez contruidos los ítems (30 a 50) serán distribuidos entre una muestra de jueces (de 50 a 100) los cuales deben ser seleccionados al azar de una población con características similares a aquella en la cual queremos aplicar la escala final. (Para los procedimientos de selección de la muestra ver el capítulo "Muestreo" en este manual.) Estos jueces responderán a cada uno de estos ítems según su opinión. Las instrucciones a los jueces pueden ser dadas según el siguiente ejemplo:

EJEMPLO DE UNA VERSIÓN PRELIMINAR

El presente es un estudio de opiniones de estudiantes universitarios respecto a problemas actuales de la universidad.

A continuación se presenta una serie de afirmaciones respecto de las cuales algunas personas están de acuerdo y otras en desacuerdo. Después de cada afirmación se ofrecen cinco alternativas de respuestas posibles:

- ☐ Totalmente de acuerdo.
- ☐ De acuerdo en general.
- ☐ Ni de acuerdo ni en desacuerdo.
- ☐ En desacuerdo en general.
- ☐ Totalmente en desacuerdo.

Indique, por favor —marcando con una cruz en el paréntesis—, *la alternativa que más se asemeje a su opinión*. Cuando no entienda alguna afirmación, ponga un signo de interrogación (?) al frente de ella. Trate de responder lo más rápido posible. Muchas gracias.

- | | |
|--|---|
| 1. Las representaciones estudiantiles deberían participar en las decisiones sobre planes de estudio. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 2. Las clases en las que el profesor tiene todo el control son las que mejores resultados producen en el aprendizaje. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 3. El plan de estudios de la UNAM debe ser centralizado en la Secretaría de Educación Pública. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 4. El trabajo en grupo es más productivo que el trabajo individual. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 5. Las carreras a las que el gobierno y la Universidad deberían prestarles más apoyo son aquellas centradas en las necesidades del país. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 6. El alumno debe tener libertad en la elección de cuál es la mejor manera de controlar su rendimiento académico. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 7. La única obligación de los alumnos es estudiar. Los planes de estudio son asunto de los profesores. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 8. La mejor manera de juzgar a un estudiante es por su rendimiento académico. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |

- | | |
|--|---|
| 9. Un trabajo hecho con consulta (en equipo o individualmente) permite una mejor evaluación de los conocimientos de los alumnos que una prueba hecha en la clase. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 10. Es preferible que cada escuela universitaria tenga un programa fijo de estudios en vez de que, como sucede en otros países, el alumno pueda escoger con alguna libertad ciertas materias de su agrado. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 11. Es preferible que los alumnos no hagan preguntas o intervenciones durante la exposición del profesor. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 12. No conviene que los alumnos intervengan en la confección de los programas de estudio. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 13. Las pruebas y exámenes deben limitarse exclusivamente a evaluar el grado de conocimiento de los alumnos respecto a la materia expuesta durante las horas de clase. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 14. Es inconveniente que ocupen plazas en la universidad estudiantes que se verán impedidos de seguir sus estudios por falta de medios económicos. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 15. Deben tener acceso a la cátedra universitaria sólo personalidades científicas de determinada orientación ideológica. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 16. No habría por qué hacer esfuerzos en adecuar los horarios de clase para la gente que trabaja. O se trabaja o se estudia. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |

- | | |
|--|---|
| 17. No deberían otorgarse becas a alumnos capaces, pero con recursos económicos suficientes, sino a aquellos de escasos recursos, aunque sean menos capaces. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 18. Sólo a los organismos centrales de dirección de la enseñanza universitaria les está dada la facultad de decir a qué ramos de la enseñanza han de conceder becas de estudios. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 19. En el organismo universitario destinado a la distribución de becas de estudio entre las distintas facultades no deben participar estudiantes. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 20. Para un mejor aprovechamiento de los recursos debe haber un organismo central que controle la concesión de becas de estudios. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 21. A la hora de tomar un acuerdo importante sobre las evaluaciones docentes universitarias, las autoridades deben hacerlo sin tener en cuenta la opinión de los alumnos. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 22. La universidad debería exigir de los egresados una retribución por los estudios recibidos, estableciendo un impuesto por el ejercicio profesional. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 23. No compete a las escuelas universitarias fijar la cantidad de matrículas anuales para sus alumnos, sino a un organismo superior central. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 24. La decisión respecto de la selección y otorgamiento a los estudiantes de material académico debe ser tomada exclusivamente por los organismos centrales de la universidad. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |

- | | |
|---|---|
| 25. La ampliación de las carreras universitarias existentes debe responder solamente al número de los postulantes que se presentan en cada escuela. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 26. En las pruebas y exámenes escritos no debe haber personal universitario que vigile a los alumnos. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 27. Los servicios de habitación y restaurante brindados por la universidad deben ser administrados sin ninguna participación del estudiantado. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 28. No debería haber un organismo central de planificación universitaria. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |
| 29. Los alumnos no deben tener injerencia alguna en la labor del personal docente. | ☐ Totalmente de acuerdo
☐ De acuerdo en general
☐ Ni de acuerdo ni en desacuerdo
☐ En desacuerdo en general
☐ Totalmente en desacuerdo |

3) Asignación de puntajes a los ítems

Con este paso comienza efectivamente el análisis de la escala. Hay que clasificar a cada ítem según sea positivo o negativo, y luego ponderar las alternativas de respuesta. Nuevamente existen diferentes criterios para la adjudicación de las ponderaciones. Por ejemplo, los pesos para un ítem *positivo* pueden ser:

Pesos

- | | |
|---|---|
| 4 | ☐ Totalmente de acuerdo |
| 3 | ☐ De acuerdo en general |
| 2 | ☐ Ni de acuerdo ni en desacuerdo |
| 1 | ☐ En desacuerdo en general |
| 0 | ☐ Totalmente en desacuerdo |

o la alternativa:

<i>Pesos</i>		
2	()	Totalmente de acuerdo
1	()	De acuerdo en general
0	()	Ni de acuerdo ni en desacuerdo
- 1	()	En desacuerdo en general
- 2	()	Totalmente en desacuerdo

o cualquier otra serie de números.

Por lo general, desaconsejamos la utilización de signos positivos y negativos en la adjudicación de los puntajes o de pesos a las alternativas de respuestas, ya que pueden crear la falsa impresión, que la escala está midiendo a nivel intervalar; esto es, donde tendríamos puntajes finales en los que existe una posición 0, posiciones +1, +2..., +40; y posiciones -1, -2, -3..., -40. De hecho la escala mide a nivel ordinal y los valores de escala simplemente implican posiciones de rango.

Para los ítems negativos, hay que recordar que la serie de números a adjudicar debe ser inversa. Por ejemplo, en un ítem negativo, la ponderación siguiendo la primera alternativa se haría:

<i>Pesos</i>		
0	()	Totalmente de acuerdo
1	()	De acuerdo en general
2	()	Ni de acuerdo ni en desacuerdo
3	()	En desacuerdo en general
4	()	Totalmente en desacuerdo

Los ítems se ubican, ya sea en forma positiva o negativa, en relación a la variable con el fin de controlar los efectos del “*response-set*”, esto es, controlar las pautas de respuesta de aquellos respondientes que tienden a dar respuestas afirmativas o negativas de manera automática.

4) Asignación de puntajes totales

Este paso consiste simplemente en la adjudicación de los puntajes totales para cada individuo en la muestra de jueces. Esta suma resulta de la adición de los puntajes ponderados para cada ítem. En el caso de ítems con valores negativos, la suma es algebraica.

De comienzo estamos asumiendo que las personas con alto grado en la variable van a tener puntajes altos, mientras que las personas con una baja actitud manifestarán puntajes bajos. Si hemos presentado 30 ítems con un valor ponderado máximo de 4, y un mínimo de cero, la amplitud total de la dispersión de la variable a esperar sería entonces

120 (puntuajes máximos de 120 y mínimo de 0, respectivamente).

5) *Análisis de los ítems*

Una vez computados los puntuajes totales para todos los jueces, hay que ordenarlos de manera que el sujeto con el puntuaje total más alto ocupe el primer lugar, el segundo puntuaje más elevado a continuación, etc., hasta llegar a la persona con el puntuaje más bajo.

Una vez ordenados los sujetos, vamos a operar únicamente con los cuartiles superiores e inferiores, es decir, el 25% de los sujetos con puntuajes más elevados y el 25% de los sujetos con puntuajes más bajos. Del 50% del centro no nos vamos a preocupar más. Formamos de esta manera un grupo *alto* y un grupo *bajo* con respecto a la variable y a los puntuajes totales. Si tuviéramos 52 jueces, el grupo alto estará constituido por los 13 jueces con los puntuajes más elevados y el grupo bajo por los 13 con puntuajes más bajos.

Tomamos a estos 26 sujetos y los colocamos en una tabla en donde situamos las puntuaciones en cada ítem y el puntuaje total para cada uno de los sujetos ordenados.

Hay que seleccionar ahora los ítems que discriminen mejor. Hay tres técnicas más en uso para la selección de los ítems: la del *cálculo del poder discriminatorio* de cada ítem; la de correlación ítem-test, y el test *de la mediana*.

El procedimiento para el cálculo del poder discriminatorio de un ítem, sigue la forma mencionada en el cuadro 3. Una vez separados el grupo alto y el grupo bajo, se calculan los promedios de cada ítem en cada uno de los grupos. Siguiendo el ejemplo de la figura 1, el promedio del ítem 1 es de 3.7 para el grupo alto (resultado de 48/13); y de 0.9 para el grupo bajo (resultado de 12/13).⁸

Una vez calculados los valores promedios para cada ítem en los grupos alto y bajo, procedemos a calcular el poder discriminativo de cada ítem según la fórmula:

$$t = \frac{DM}{\sqrt{\frac{s_{M_1}^2}{N_1 - 1} + \frac{s_{M_2}^2}{N_2 - 1}}}$$

Donde:

t = *Test t* de Student

DM = Diferencia entre medidas ($M_1 - M_2$)

s^2 = variancias de muestra 1 y 2 respectivamente

N = cantidad de casos en cada una de las muestras

Para el cálculo del poder discriminativo del ítem conviene utilizar el siguiente cuadro, que es una continuación del cuadro 3:

Ordenamiento de jueces	Ítem núm.																														Puntaje total
1	4	4	4	4	3	4	2	4	3	4	4	4	etc...																		112
2	4				2		2																								108
3	3				4		0																								108
4	4				0		1																								104
5	4				1		2																								102
6	4				0		0																								101
7	4				0		0																								100
8	3				4		2																								98
9	4				2		1																								98
10	3				3		1																								96
11	3				3		1																								97
12	4				2		0																								97
13	4				2		2																								96
40	2				2		3																								27
41	0				3		4																								26
42	2				0		2																								26
43	1				1		4																								26
44	0				4		4																								24
45	1				0		3																								17
46	0				1		3																								16
47	1				2		4																								16
48	2				1		4																								14
49	3				1		3																								14
50	0				2		4																								8
51	0				3		4																								8
52	0				4		3																								7

CUADRO 4. Valores promedios para los grupos alto y bajo y diferencia de medias, para el cálculo del poder discriminatorio de ítems en la versión de los jueces en una escala Lickert.

	1	2	3	4	5	6	7	8	9	10	11	12
Promedio en el grupo alto (M_1)	3.7				1.8		1.1					
Promedio en el grupo bajo (M_2)	0.9				1.7		3.4					
Diferencia de medias ($M_1 - M_2$)	2.8				0.1		-2.3					

Una vez calculados los valores t , confrontamos con la tabla de distribución de valores correspondientes, seleccionando aquellos ítems que realmente presenten diferencias significativas entre ambos grupos de contraste.

En los ejemplos señalados en la figura 1, la diferencia para el ítem 1 es significativa al nivel de 0.1 y la del ítem 7 al nivel de 0.5. El ítem 5 debe ser desechado porque no discrimina significativamente. Es importante notar asimismo que en el caso del ítem 7

hemos colocado el signo del ítem originalmente mal ya que, como vemos, el grupo con valores bajos responde al estímulo en forma más positiva que los del grupo alto. Corresponde entonces cambiar los valores, manteniendo el ítem.

El método *ítem-test* para la selección de los ítems consiste en correlacionar el puntaje del ítem con el puntaje total del test. El coeficiente de correlación a utilizar es el coeficiente de correlación biserial, ya que aquí se trata de dos variables, una que podemos considerar intervalar, y la otra una serie dicotomizada, siendo además la dicotomía forzada. El supuesto que tenemos que aceptar es que la distribución original es continua y normal. La fórmula para el cálculo de la correlación será:

$$r_b = \frac{M_p - M_q}{\sigma_t} \cdot \frac{p \cdot q}{y}$$

Donde:

M_p y M_q = Medias parciales en el grupo alto y bajo, respectivamente.

σ_t = Desviación estándar total.

p y q = Proporción de casos en una y otra distribución.

y = Ordenada a la curva normal.

Cualquiera que sea la técnica utilizada para el análisis de los ítems, el objetivo es seleccionar aquellos que discriminen mejor (valores significativos de t , o valores altos de r_b).

De hecho hay una alternativa al uso de la prueba t , o del coeficiente de correlación biserial, que estrictamente se corresponde más con el tipo de medición de nivel ordinal, con el que opera la escala Lickert. En este caso la diferencia entre medianas, computada a través del “test de la mediana”. Para ello hay que determinar, primero, el valor de la mediana de cada ítem para los grupos alto y bajo combinados, luego dicotomizamos los valores en una tabla de 2×2 para cada ítem, de la siguiente forma

	Grupo alto	Grupo bajo	
Número de puntajes por debajo de la mediana combinada	A	B	A + B
Número de puntajes por encima de la mediana combinada	C	D	C + D
	A + C	B + D	

A esta tabla aplicamos ya sea χ^2 o el *test* de Fisher, según sea la cantidad de casos (por lo menos $N = 40$ para aplicar χ^2 ; con menos casos se recomienda el *test* de Fisher).

La fórmula para el cómputo de χ^2 para este tipo de figura es:

$$\chi^2 = \frac{N \left[(AD - BC) - \frac{N}{2} \right]^2}{(A + B)(C + D)(A + C)(B + D)}$$

Seleccionamos por supuesto aquellos ítems cuyo χ^2 da diferencias significativas. Usando la misma figura, el test de Fisher se computa según la siguiente fórmula:

$$p = \frac{(A + B)! (C + D)! (A + C)! (B + D)!}{N! A! B! C! D!}$$

6) La versión final de la escala

La cantidad de ítems seleccionados de acuerdo a su poder discriminativo constituye la escala final por aplicar a sujetos o grupos como versión final. Los puntajes finales por adjudicar a los individuos son entonces el producto de la suma de los puntajes obtenidos en cada ítem, divididos entre el total de ítems.

B) Comentarios finales

Para calcular la *confiabilidad de la escala*, se puede utilizar la correlación entre mitades del test (*split-half reliability*): se correlacionan la suma de los puntajes en los ítems impares con la suma de los puntajes de los ítems pares. Utilizamos ρ de Spearman, y luego la fórmula:

$$C = \frac{2\rho}{1 + \rho}$$

A continuación presentamos las ventajas y desventajas de la escala Lickert, comparada con la Thurstone.

Ventajas. a) Permite la utilización de ítems que no se encuentran relacionados en forma manifiesta con la actitud que se desea estudiar (es decir, se pueden utilizar ítems con contenido latente). b) Es más rápida y fácil de construir. c) A mismo número de ítems, es más confiable. d) La cantidad de alternativas de respuesta permite una información más precisa de un sujeto en un ítem particular.

Desventajas. a) Por tratarse de una escala ordinal, no permite apreciar la distancia que hay entre pares de sujetos con respecto a la actitud. b) Con frecuencia dos puntajes iguales pueden ocultar pautas de respuestas diferentes de los individuos. c) No hay garantía de unidimensionalidad, consecuentemente pueden mezclarse distintas dimensiones, no estando seguro el investigador de cuál de ellas realmente se trata.

Thurstone es quien provee la racionalidad —mediante su *ley de juicios comparativos*— para todo el aparato conceptual en la construcción de escalas para medir actitudes. Esta ley sostiene que, para cada estímulo (e) dado, está asociado un proceso modal discriminacional sobre un continuo psicológico. La distribución de todos estos procesos discriminacionales sigue la forma de la distribución normal, en la que todos los procesos discriminacionales producidos por el estímulo se distribuyen normalmente alrededor del proceso de discriminación modal, con una dispersión discriminacional (s_i). Dado un conjunto n de estímulo, es posible ordenarlos en un continuo psicológico tomando como referencia el grado de atributo que ellos poseen.

A partir de estos principios, Thurstone propone su *escala de intervalos aparentemente iguales*, de tipo diferencial, en la que los ítems son seleccionados por una serie de técnicas que permiten escalonarlos de manera tal que expresen el continuo psicológico subyacente. La medición trata de establecerse al nivel intervalar. Es decir, una escala en la que sea posible afirmar que la distancia que separa a un sujeto que obtuvo una puntuación de 8.7 con respecto a otro sujeto que obtuvo 6.3 es igual a la distancia que separa a otro par de sujetos que obtuvieron puntuaciones de 3.6 y 1.2, respectivamente, y a cualquier distancia que sea igual a 2.4 puntos. (Sin embargo, como veremos más adelante, es discutido que la escala mida efectivamente este nivel.)

El continuo psicológico en la escala de intervalos aparentemente iguales de Thurstone se edifica sobre una serie de juicios de actitud distribuidos en una escala de 11 puntos, en la que el punto 1 de la escala representa una actitud extrema (favorable o desfavorable), el punto 6 representa una actitud neutra (ni favorable ni desfavorable), y el punto, 11 el otro extremo (favorable o desfavorable, según el extremo contrario a la actitud asumida en 1).

Los ítems en la escala Thurstone son contruidos, diseñados y seleccionados de manera tal que permitan atribuir a los sujetos a los que se aplicará definitivamente la escala un punto en un continuo. Así, esta escala es un poco más refinada que la escala Lickert, e implica una cantidad considerable de trabajo adicional.

A) La construcción de una escala Thurstone

Los procedimientos para la construcción de una escala de este tipo son: 1) Se construye una serie de ítems (alrededor de 150). 2) Se solicita a un grupo de jueces (más o menos 100) que ubiquen los ítems en una escala de 11 puntos. 3) Una vez evaluados los ítems por los jueces, se adjudica a los ítems valores de escala. 4) Se seleccionan los ítems que representan el rango entero de la escala, rechazando los ítems ambiguos.

Detallamos cada uno de los pasos:

1) La construcción de los ítems (versión de los jueces)

Construya entre 100 y 200 ítems tomando en consideración los siguientes criterios (ver

Edwards, A. L., para mayores detalles):

- a) Evite los ítems que señalan el pasado en vez del presente.
- b) Evite los ítems que dan demasiada información sobre hechos o los que fácilmente puedan ser interpretados como tales.
- c) Evite los ítems ambiguos.
- d) Evite los ítems irrelevantes con respecto a las actitudes que pretende medir.
- e) Evite los ítems con los cuales todos o nadie concuerda.
- f) Los ítems deben ser formulados en un lenguaje simple, claro y directo.
- g) Sólo en casos excepcionales sobrepase las 20 palabras en un ítem.
- h) Un ítem sólo debe contener una frase lógica.
- i) Los ítems que incluyan palabras como “todos”, “siempre”, “nadie”, “nunca”, etc. serán percibidos de la misma manera y por ello deben omitirse.
- j) De ser posible, los ítems deben ser formulados como frases simples y no compuestas.
- k) Use sólo palabras que el entrevistado pueda comprender.
- l) Evite las negaciones, especialmente las dobles negaciones.
- m) Combine los ítems formulados positiva, neutral y negativamente en una proporción de 1/3, 1/3 y 1/3, distribuidos uniformemente sobre la variable.

Algunas maneras de formular ítems pueden ser: *i)* Extraerlos de libros, publicaciones y artículos que tratan sobre el objeto cuya actitud se quiere medir. Son importantes también las declaraciones y los discursos, a partir de los cuales uno pueda hacer una especie de análisis de contenido. Cuando se sigue esta estrategia hay que tener en cuenta que un buen monto de reformulación va a ser necesario. *ii)* Concertar una discusión entre personas que representan distintos puntos de vista con respecto a la actitud. En este caso la grabación de la discusión facilitará la selección de frases adecuadas. Aquí también será necesaria la reformulación. *iii)* Formular uno mismo, o en cooperación con otros investigadores, los enunciados ante los cuales se espera que la gente reaccione en forma positiva, negativa o neutra.

En todo caso nunca es fácil llegar a la enunciación de 100-200 ítems sin incurrir en repeticiones o en formulaciones muy similares. Es importante, por cierto, que la distribución de los ítems a presentar a los jueces sea aproximadamente pareja, conteniendo un tercio de ítems positivos, negativos y neutros, a lo largo de un continuo.

2) La administración a los jueces

La lista de ítems (100 a 200) se distribuye a jueces (preferiblemente 200, con un mínimo de 50), los cuales van a ubicar los ítems en una escala intervalo-subjetiva que va de 1 a 11 puntos.

Los jueces son seleccionados en función de su *conocimiento* sobre el problema que se quiera medir; y en la clasificación o evaluación de los ítems no importa la opinión personal del juez, sino su evaluación del punto en la escala continua de 1 a 11 en el cual él ubica el ítem (es decir, la determinación del peso que el ítem tiene en su opinión para

la medición de la actitud).

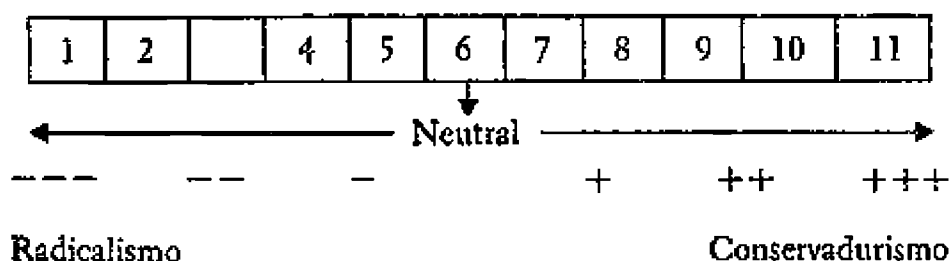
Las instrucciones a los jueces —que operan en forma independiente— pueden ser las siguientes:

EJEMPLO: ESCALA ACTITUD TIPO THURSTONE

Instrucciones

El Seminario de Recolección y Análisis de Datos del Centro de Estudios Sociológicos de El Colegio de México está realizando una serie de ejercicios sobre construcción de escalas para la medición de actitudes.

En el presente caso, tiene usted en sus manos una serie de afirmaciones para construir una escala de tipo Thurstone. La escala de intervalos aparentemente iguales de Thurstone parte de una serie de supuestos y técnicas en la cual los ítems son escalonados de manera tal que expresen un continuo subyacente. El continuo se edifica sobre una serie de juicios de actitud distribuidos en una escala de 11 puntos.



El punto 1 representa una actitud extrema (en nuestro caso: radicalismo), el punto 6 una actitud neutra (ni radical ni conservador) y el punto 11 el extremo conservador.

Para la construcción de esta escala, se requiere someter una cantidad de ítems a un número limitado de *jueces* antes de ser aplicada a una muestra de una determinada población.

Los jueces se eligen entre personas que tienen conocimientos especializados sobre la variable que se trata de medir, como ha sido el caso con usted.

En el presente caso, se trata de construir una escala que mida actitudes de radicalismo-conservadurismo. Quisiéramos pedirle que, en relación a esto, *no* vuelque usted *sus propias* opiniones acerca de las afirmaciones que aparecen a continuación, sino que usted exprese su juicio acerca de cuán radical, conservadora o neutra le parece cada una de las afirmaciones.

A la izquierda de cada afirmación hay un cuadrado en el cual usted debe colocar, de acuerdo a su criterio, el número que representa la frase en el continuo de 1 a 11, siendo:

1 el extremo "radical",

6 el punto medio o "neutral",

11 el extremo "conservador".

Si cree que la expresión se ubica *entre alguno de estos puntos*, utilice el número intermedio que mejor represente la posición de la frase. No trate de obtener el mismo número de ítems, o cualquier otra distribución espacial, en el continuo de 1 a 11.

Por favor, antes de empezar, lea una buena cantidad de expresiones para entender el carácter de los ítems.

Queremos insistir en que *no* se trata de dar una opinión personal de acuerdo o desacuerdo acerca de cada afirmación, sino solamente de estimar su lugar en una escala de 1 a 11 puntos.

[Para el ejemplo presentamos solamente los primeros 29 ítems]

- ☐ 1. El modelo desarrollista fortalece la desigualdad social.
- ☐ 2. Tratándose de programas políticos, "más vale malo conocido que bueno por conocer".
- ☐ 3. Todos los asuntos económicos de interés nacional deben estar a cargo exclusivo del Estado.
- ☐ 4. Los trabajadores deben tomar en sus manos la conducción del Estado; el Estado debe tomar en sus manos la conducción de la economía.
- ☐ 5. La desigualdad social ha existido siempre, y es necesaria para el desarrollo de la sociedad.
- ☐ 6. Los sindicatos deberían limitar sus actividades a las reivindicaciones económicas de sus representados.
- ☐ 7. La función del Estado es la de participar dinámicamente en el desarrollo económico, social y político del país.
- ☐ 8. La propiedad privada es un derecho natural del hombre y debe ser respetado y mantenido.
- ☐ 9. Las políticas sobre la distribución del ingreso son pura demagogia que sólo propician el enriquecimiento de los políticos.
- ☐ 10. La principal causa de la inflación es el anhelo de los empresarios de aumentar sus ganancias.
- ☐ 11. El Estado debería dejar absolutamente a criterio de los padres de familia el tipo de educación que prefieren para sus hijos.
- ☐ 12. Todo grupo de interés debe tener igual representación ante el Estado.
- ☐ 13. Los trabajadores producen la riqueza; los patronos se la embolsan.

- ☐ 14. La política del control de la natalidad es una política al servicio del imperialismo.
- ☐ 15. Los que más pierden con la inflación son los empresarios.
- ☐ 16. La importancia creciente del sindicalismo en el país representa un peligro para la democracia.
- ☐ 17. Antes de la Reforma Agraria se producía más y mejor en el campo.
- ☐ 18. La mejor manera de resolver los problemas es encontrar el justo “término medio” y no caer en los extremos.
- ☐ 19. No se trata del sistema tal o cual; para que el país progrese hay que trabajar más, y punto.
- ☐ 20. El camino de México es el de una economía mixta: el gobierno como promotor, y la iniciativa privada como participante activo en el proceso de desarrollo.
- ☐ 21. El éxito se debe al esfuerzo personal.
- ☐ 22. La nacionalización de la industria minera sólo caería en el burocratismo y la mala administración estatal.
- ☐ 23. La mejor forma de representación ante el Estado es a través de grupos de interés, y no a través de partidos políticos.
- ☐ 24. La intervención del gobierno en la economía agrícola sólo ha traído desorden y caos.
- ☐ 25. El deficiente desarrollo agrícola del país se debe a la apatía y flojera de los campesinos.
- ☐ 26. La Revolución mexicana será verdadera sólo cuando se realice una reforma agraria total.
- ☐ 27. La fuerza política de un grupo debe ser independiente del poder económico de sus miembros.
- ☐ 28. La actividad económica es privativa de los particulares y el Estado debe limitarse a coordinar tal actividad.
- ☐ 29. El marxismo es una doctrina exótica que no toma en cuenta nuestra idiosincrasia ni nuestra tradición.

3) Asignación de valores de escala

Vamos a estudiar ahora la distribución de las respuestas (de 1 a 11) en cada ítem, según las respuestas sobre ubicación del ítem dadas por los jueces. Podríamos calcular los valores promedios (media aritmética) y las desviaciones estándar (σ), para cada ítem. Sin embargo —y aquí de hecho se revela que la escala es más ordinal que intervalar—, existe mayor exactitud y representa un método más rápido el calcular valores de mediana (Mdn) y de distancia intercuartil (Q_3-Q_1). En las siguientes páginas ilustramos con

ejemplos de análisis gráficos. (Ver cuadros 5, 6 y 7.)

En el primer ejemplo (*cuadro 5*) hemos obtenido una mediana de 2.3. La mediana indica el “valor” del ítem a lo largo de la variable, es decir que se trata en este caso de un ítem al lado positivo de la escala (muy dogmático). La distancia intercuartil es igual a 2.2. Esta cifra indica la “calidad” del ítem. Si menor la distancia intercuartil, mayor el grado de “calidad” del ítem, es decir que la adjudicación del valor del ítem por parte de los jueces es similar. Los valores altos de distancia intercuartil, por el contrario, indican diferencias entre los jueces en cuanto a la apreciación sobre el valor adjudicado al ítem.

El ejemplo 2 (*cuadro 6*) indica el caso ideal, en el que todos los jueces están de acuerdo en cuanto al valor que debe ser adjudicado al ítem.

El ejemplo 3 (*cuadro 7*) representa un ítem de escasa utilidad, ya que como se ve los jueces le adjudicaron valores muy distintos, representando el valor mediano del ítem una adjudicación casi aleatoria de los jueces en cuanto al valor 1, 2, 3... u 11 que se le adjudique al ítem como expresando una posición en el continuo.

En los tres ejemplos presentados se han utilizado 56 jueces, y el método para calcular la mediana y los cuartiles ha sido proyectando líneas sobre la distribución gráfica. A partir de la distribución de frecuencias sobre los valores de ítems (1 a 11), resulta simple calcular la mediana y los cuartiles según las siguientes fórmulas:

$$Mdn = v + i \left(\frac{\frac{N}{2} - F_d}{F_p} \right)$$

Donde:

v_i = Límite exacto inferior del intervalo que contiene la mediana.

i = Amplitud del intervalo de clase (en nuestro caso igual a 1).

N = Número de casos.

F_d = Suma total de las frecuencias inferiores al intervalo que contiene la mediana.

F_p = Frecuencias del intervalo que contiene la mediana.

Para el cálculo de cuartiles:

$$Q_3 = v_s + i \left(\frac{\frac{N}{4} - F_d}{F_p} \right)$$

$$Q_1 = v_i + i \left(\frac{\frac{N}{4} - F_d}{F_p} \right)$$

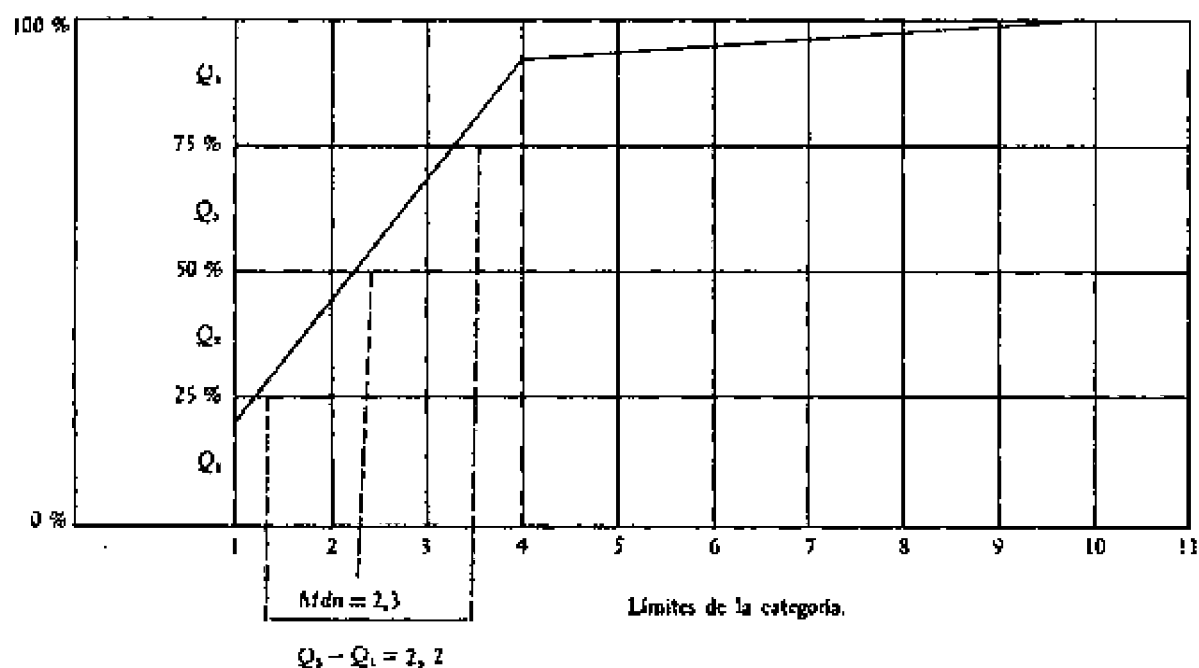
CUADRO 5. Ejemplo 1: Escala de actitud ítem bueno tipo Thurstone mide dogmatismo

Ítem núm. 173

Caja núm.	1	2	3	4	5	6	7	8	9	10	11
Frecuencia	//// //// //	//// //// /	//// //// //	//// //// ///	////	/	//				/
Frecuencia en números	12	11	12	13	4	1	2				1
Frecuencia acumulativa	12	23	35	48	52	53	55				56
Porcentaje acumulativo	21.4	41.1	62.5	85.7	92.5	94.6	98.3				100
	1	2	3	4	5	6	7	8	9	10	11

56 jueces

Representación gráfica y cálculo de Mdn y Q



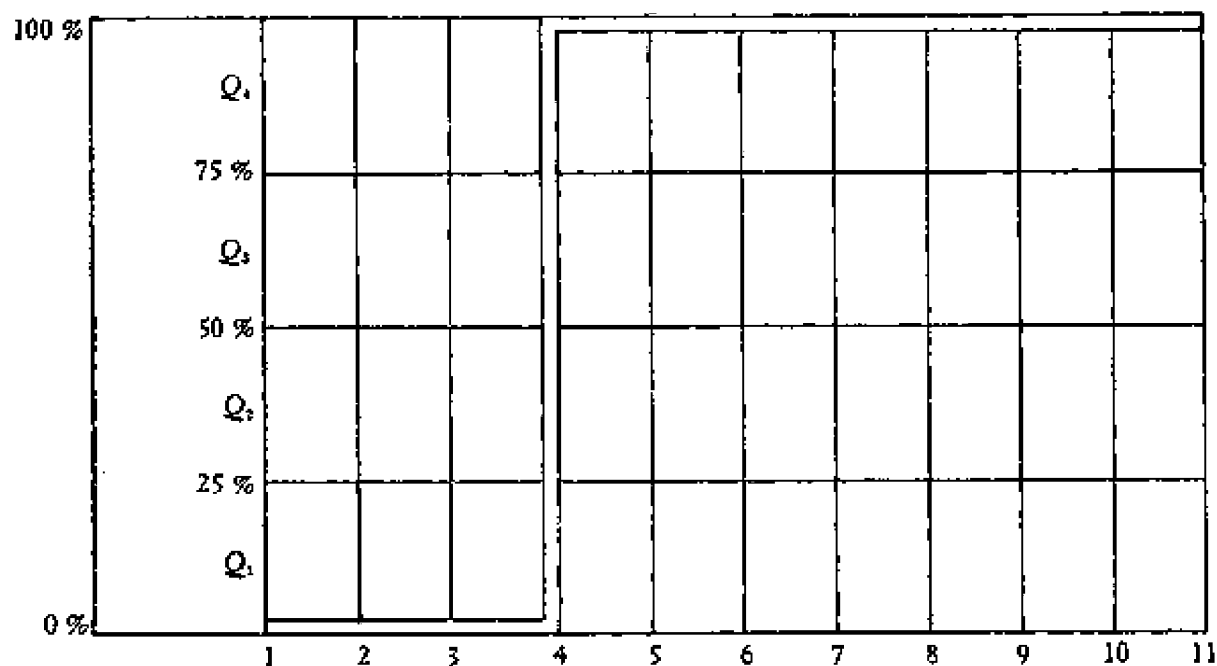
CUADRO 6. Ejemplo 2. Escala de actitudes tipo Thurstone, ítem ideal

Ítem núm. X

Caja núm.	1	2	3	4	5	6	7	8	9	10	11
Frecuencia											
Frecuencia en números	0	0	0	56	0	0	0	0	0	0	0
Frecuencia acumulativa	0	0	0	56	56	56	56	56	56	56	56
Porcentaje acumulativo	0	0	0	100	100	100	100	100	100	100	100
	1	2	3	4	5	6	7	8	9	10	11

56 jueces

Representación gráfica y cálculo de Mdn y A



$$Mdn = 4$$

Límites de categoría.

$$Q_3 - Q_1 = 0$$

CUADRO 7.

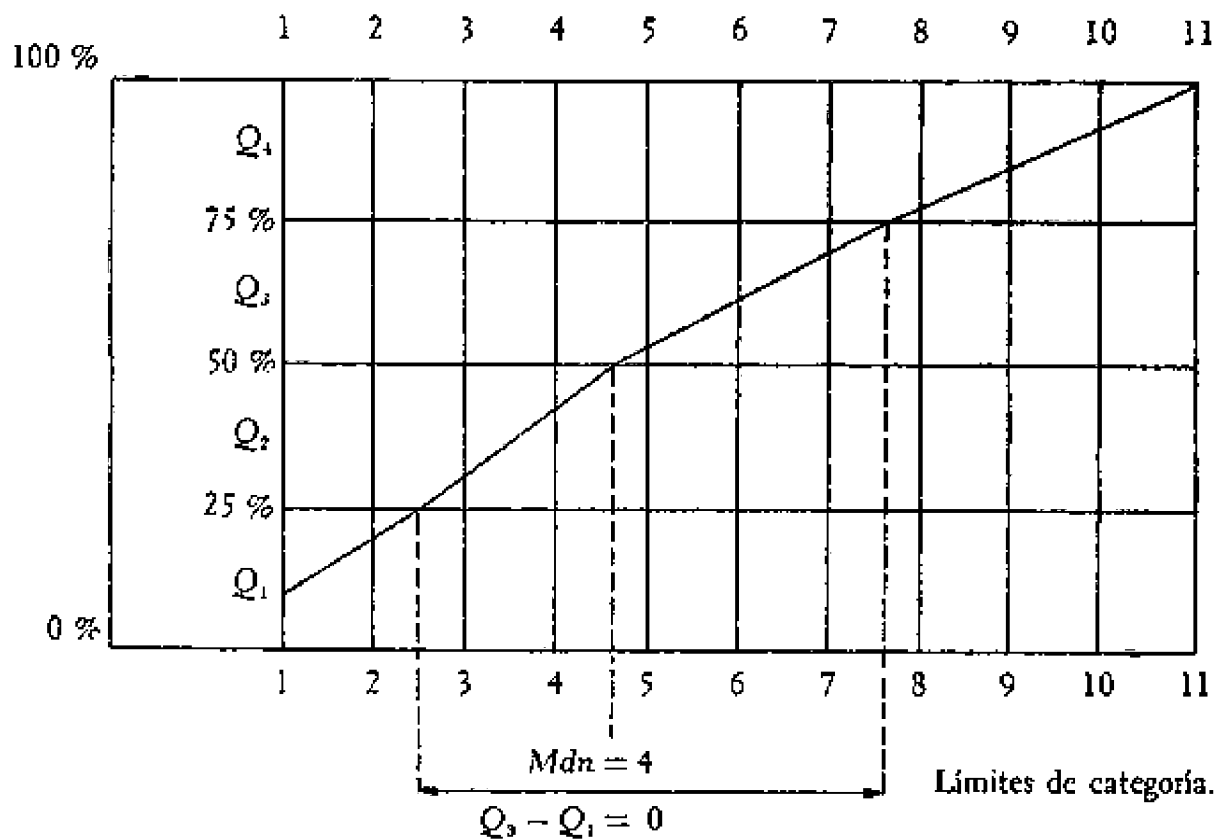
Item núm.

172

Caja núm.	1	2	3	4	5	6	7	8	9	10	11
Frecuencia											
Frecuencia en números	3	7	8	6	6	4	6	4	5	2	5
Frecuencia acumulativa	3	10	18	24	30	34	40	44	49	51	56
Porcentaje acumulativo	5.4	17.8	32.2	42.9	53.6	60.8	71.2	78.5	87.5	91.0	100
	1	2	3	4	5	6	7	8	9	10	11

56 jueces

Representación gráfica y cálculo de Mdn y A.



Donde:

v_s = Límite exacto superior del intervalo que contiene el cuartil.

v_i = Límite exacto inferior del intervalo que contiene el cuartil.

i = Amplitud del intervalo de clase.

F_s = Frecuencias por encima del intervalo que contiene el cuartil.

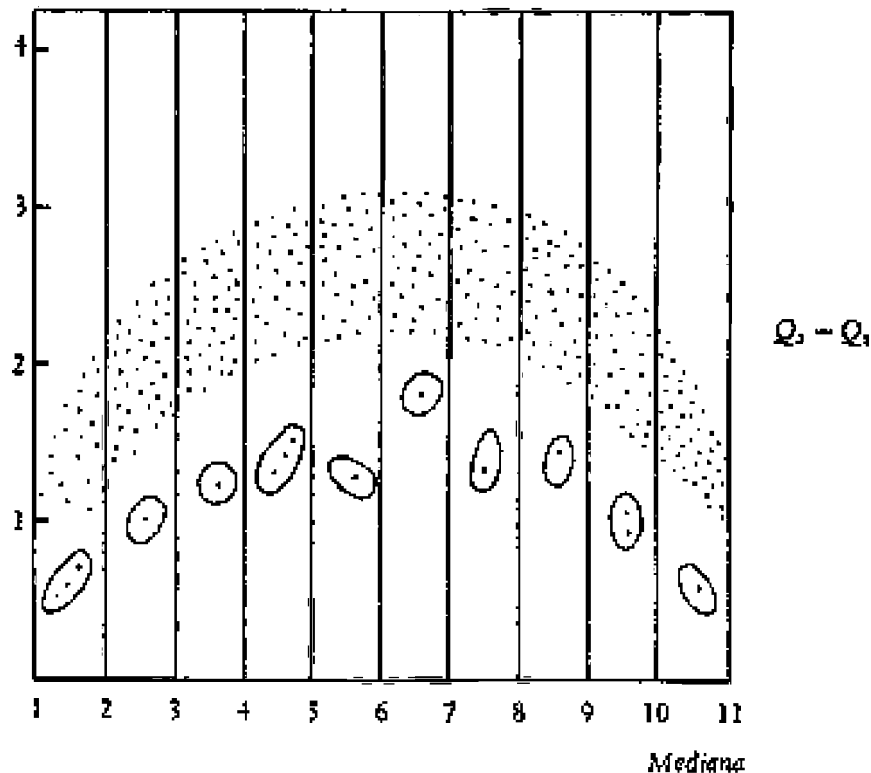
F_d = Suma total de frecuencias por debajo del intervalo que contiene el cuartil.

F_p = Frecuencias en el intervalo que contiene el cuartil.

4) La selección de los ítems

La selección final de los ítems presentados a los jueces se realiza en base a los valores de la mediana, y de amplitud intercuartil. La mediana se utiliza para ubicar el peso del ítem en la escala, eligiendo los ítems que se hallen repartidos uniformemente a lo largo de la misma. Necesitamos aproximadamente 2 ítems para cada intervalo de la escala (1-2; 2-3; 3-4 ... 9-10; 10-11).

FIGURA 3.



La distancia intercuartil es utilizada para decidir cuáles son los mejores ítems dentro de cada intervalo. La figura 3 indica una distribución de ítems en la cual en el eje horizontal figuran los valores de mediana, y en el eje vertical los valores de distancia intercuartil. Los ítems seleccionados para la escala final están marcados con círculos y son

obviamente aquellos que, representando cada uno de los intervalos en la escala, presentan distancias intercuartiles mínimas.

Normalmente las distancias intercuartiles más grandes se presentan en el centro de la escala. Ocurre muchas veces que la selección de ítems para los valores 5, 6, 7 y 8 se complica en la medida en que los jueces tienden a dispersar sus evaluaciones fuertemente en el centro; es decir, esto se va a reflejar en distancias intercuartiles bastante grandes. Si no se obtienen suficientes ítems en la primera versión de los jueces, el investigador deberá construir más ítems para esta área, que deben ser juzgados nuevamente por los mismos jueces. Conviene por supuesto anticiparse a este tipo de problemas, teniendo especial cuidado cuando se construyen los ítems iniciales.

B) La versión final de la escala

Está constituida por los 15 o 25 ítems seleccionados como confiables. Los ítems son presentados a los sujetos solamente con dos alternativas de respuestas: acuerdo-desacuerdo.

La forma de presentación de la escala es diferente de aquella presentada a los jueces. En un cuestionario, por ejemplo, se eliminan las “cajas”, solicitando al sujeto que responda únicamente si está de acuerdo o en desacuerdo con cada uno de los juicios.

A manera de ejemplo, presentamos algunos ítems que han sido contestados por un sujeto hipotético.

1. El éxito se debe al esfuerzo personal.	☒ De acuerdo
	☐ En desacuerdo
2. El futuro nos deparará mejores condiciones de vida.	☐ De acuerdo
	☒ En desacuerdo
3. El gobierno debería tomar todas las decisiones.	☐ De acuerdo
	☒ En desacuerdo
4. Casi nadie considera el trabajo que yo hago.	☐ De acuerdo
	☒ En desacuerdo
5. Es bueno que la Iglesia se modernice.	☐ De acuerdo
	☒ En desacuerdo
6. Yo no podría ser bueno de carácter.	☒ De acuerdo
	☐ En desacuerdo

Etcétera.

5) La adjudicación de puntajes a los sujetos

El investigador tiene que registrar en un código los valores de mediana adjudicados por los jueces a los ítems seleccionados para integrar la versión final de la escala. Los valores intercuartiles no se utilizan para el cómputo de los puntajes finales de los sujetos. El puntaje final de un sujeto será entonces simplemente el promedio de los valores de escala de los ítems respondidos en forma afirmativa o “de acuerdo”.

Supongamos que nuestro sujeto hipotético haya respondido “de acuerdo” únicamente a los ítems 1 y 6; es decir que en el resto de los ítems su respuesta ha sido “en desacuerdo”. Supongamos entonces que los valores de escala dados por los jueces hayan sido:

<i>Ítem</i>	<i>Peso del ítem en la variable</i>
1	1.2
2	4.6
3	7.8
4	3.7
5	10.8
6	1.6
7	9.8
8	5.6
etc.	etc.

Consecuentemente, el puntaje correspondiente a nuestro sujeto será:

$$\frac{1.2 + 1.6}{2} = 1.4$$

Los valores en la escala utilizados como una variable en el análisis de alguna investigación pueden también ser utilizados de varias formas. Se pueden tomar los distintos valores en los sujetos y establecer correlaciones entre la actitud y alguna otra variable, o bien dicotomizar las respuestas de la siguiente manera:

<i>Peso en la escala Thurstone</i>	<i>Núm. de individuos</i>	<i>Frecuencias acumulativas</i>	
1	34	34	
2	54	88	
3	42	130	
4	39	169	
5	15	184	grupo
6	5	189	alto
7	52	241	
8 Dicotomización	84	325	
9	82		grupo
10	72		bajo
11	121		
	600		

Supongamos que la distribución en nuestra variable entre 600 individuos en una encuesta es como aparece en la página anterior (supongamos que todos han contestado la pregunta).

Lo que hicimos con la distribución fue lo siguiente:

a) Contabilizamos el número de sujetos que obtuvieron los diferentes puntajes en la escala; aquí particularmente ubicamos a los sujetos que obtuvieron, por ejemplo, 4.6 en el intervalo 5. Nótese que, entonces, la amplitud de cada intervalo de “peso en la escala”, para el caso del intervalo 5, va de 4.5 a 5.49.

b) Calculamos las frecuencias acumuladas,

c) Si buscamos dicotomizar la variable, el estadístico a utilizar sería la mediana, que en nuestro caso tiene un valor igual a 8.2, es decir, cae en la categoría 7.5 a 8.49.

d) Procedemos entonces a la dicotomización y obtenemos 2 grupos, a los cuales les llamaremos grupo *alto* y grupo *bajo*.

e) Podemos cruzar entonces nuestra variable con alguna otra, tal como figura en el ejemplo,

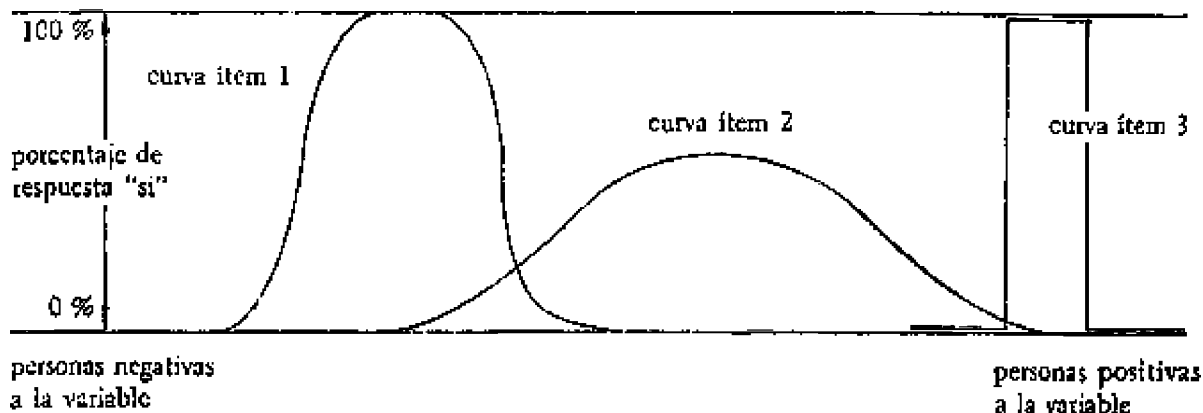
		Variable Y (otra variable cualquiera)			
		Alto	Mediano	Bajo	
Variable X (Thurstone)	Grupo alto	15	130	180	325
	Grupo bajo	221	52	2	275
		236	182	182	600

C) Comentarios

1) Ítems o series de ítems acumulativos o diferenciados

Los ítems de la escala Thurstone son ítems diferenciadores. Presentamos algunos ejemplos de diferentes ítems en una escala Thurstone. (Ver en la página siguiente.)

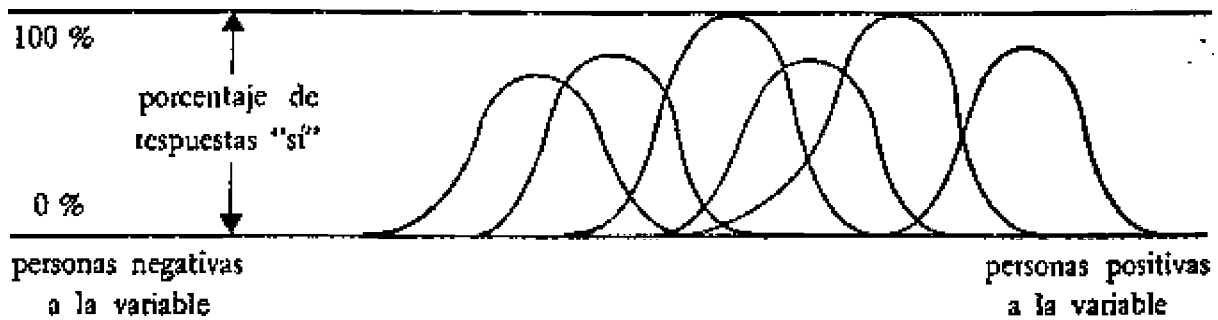
El eje horizontal representa todas las personas investigadas ordenadas en tal forma que las que detentan una actitud más negativa con respecto a la variable se ubican a la izquierda y las más positivas a la derecha. En el eje vertical está representado el porcentaje de respuestas “de acuerdo” en un ítem.



ítem 1: es un ítem que recibe respuestas “No” de las personas que tienen una fuerte posición ya sea positiva o negativa. Un grupo de personas que tengan una determinada posición a lo largo de la escala responderán este ítem con un “Sí”. En el caso del ítem 2, la única diferencia es que solamente el 50% de las personas responderá “Sí” al ítem en el lugar de la escala donde la probabilidad de alta tasa de respuestas “Sí” es más grande. El ítem 3 es un ítem donde el 100% de los respondientes con ciertas posiciones en la escala responderán “Sí” y todos los demás “No”. Ejemplo: “¿Está usted entre los 178 y los 179 cm de estatura?”

La escala Thurstone en su totalidad forma una serie de ítems diferenciadores. La

siguiente figura puede dar una idea de la estructura de una escala Thurstone. Hemos reducido el número de ítems a 6 para facilitar la lectura de la figura.



2) La variación en el instrumento, los sujetos o ambos

En la escala de Thurstone, solamente el estímulo aislado recibe valor en la escala. El conjunto de tareas para los jueces es escoger los ítems dentro de los intervalos de igualdad-apariencia (11 cajas) tratando de disminuir cualquier variación debido a su *propia* posición con respecto a la actitud. Es decir que en el caso de la escala Thurstone tenemos el enfoque que corresponde a la variación centrado en el instrumento. La escala Thurstone entonces pertenece al test Tipo B de Coombs.

3) El problema con las dimensiones

Una desventaja de la escala de Thurstone es la dificultad de controlar las dimensiones en la variable. Supongamos que una persona ha contestado “de acuerdo” a los siguientes ítems de una escala Thurstone de 30 ítems:

<i>Item</i>	<i>Valor</i>
4	5.9
9	6.2
11	6.9
19	7.2
28	7.8
	34.0

Su promedio es entonces $\frac{34.0}{5} = 6.8$

Pudo haber recibido esa suma particular mediante una serie de operaciones. Los cinco ítems con los que concordó pudieron haber tenido asignados los números 4.1; 4.2; 8.7; 8.5; 8.9 o por ejemplo: 1.9; 2.2; 7.4; 10.7 y 10.8. En ambos casos el promedio es 6.8. Las posibilidades aumentan todavía más si consideramos que, de hecho, el individuo podía haber concordado con más o menos de 5 ítems. ¿Cuál podría ser la razón de esto? Como

simplificación gruesa del problema podemos pensar que la serie 4.1; 4.2; 8.7; 8.5 y 8.9 indica dos dimensiones, una de las cuales aparece alrededor de 4 y la otra alrededor de 9. Aparentemente, los dos primeros ítems tienen para el sujeto un significado distinto de los tres siguientes. Para controlar esta desventaja de la escala, Thurstone y Chave (1929) proveyeron a la escala de un criterio burdo basado en las respuestas: el llamado *index of similarity* (índice de similitud).

D) Ventajas y desventajas de la escala Thurstone

Ventajas. a) La principal ventaja de esta técnica es que permite hacer una distribución de un grupo dado, a lo largo de la actitud que se desea investigar, precisamente porque los ítems fueron diseñados y seleccionados a los efectos de cubrir el continuo. b) Supone una medida más refinada que la escala Lickert, ya que el puntaje de los ítems se deriva de una ponderación basada en el juicio de jueces expertos o al menos informados. c) En la medida en que la escala final contiene más ítems que la de Lickert, es más confiable que ésta. d) Si es tratada como escala intervalar, permite comparar puntuaciones y cambios de actitud en los sujetos.

Desventajas. a) Su elaboración es larga y compleja. b) Pese a su tratamiento como escala intervalar, su verdadera naturaleza en términos de nivel de medición corresponde al nivel ordinal. c) Fácilmente se introducen otras dimensiones distintas de las que se quiere medir. d) Discrimina poco en los extremos de la distribución. e) Los ítems neutrales carecen de significado mayor, y a menudo se ubican en esta posición ítems que no se refieren a la dimensión tratada. f) Distintas configuraciones de respuesta resultan en el mismo puntaje final. g) Los jueces pueden introducir sesgos difíciles de detectar.

LA ESCALA GUTTMAN

Una de las desventajas mayores en las dos escalas que examinamos hasta ahora — Lickert y Thurstone— era que ninguna de ellas garantiza que el instrumento mida una dimensión única.

La escala Guttman, conocida como método del escalograma o análisis de escalograma, soluciona el problema de la unidimensionalidad. Su objetivo es definir lo más claramente posible qué es lo que está midiendo la escala, entendido esto como un problema de unidimensionalidad. Por el tipo especial de tratamiento al que se somete la escala se busca la eliminación de factores extraños a la característica o dimensión que se pretende medir.

La escala Guttman es de tipo acumulativo, ya que la respuesta positiva a un ítem supone que los ítems anteriores también han sido respondidos en forma positiva. Se busca pues una coherencia en las pautas de respuesta de los sujetos, y esa coherencia es garantizada por medio de un *coeficiente de reproducibilidad*. El tamaño del coeficiente (valor máximo 1.00) señala el grado por el cual la escala es acumulativamente perfecta o casi perfecta. En una escala cuya reproductividad es perfecta, las respuestas de los sujetos a todos los ítems pueden ser reproducidas por el solo conocimiento de su

posición de rango.

Veremos más adelante que, además de la reproductividad, hay que tomar en cuenta otros factores, tales como el alcance de la distribución marginal, la pauta de los errores, el número de ítems en la escala y el número de categorías de respuesta.

El método de escalograma de Guttman combina aspectos de construcción utilizados en las escalas Lickert y Thurstone, además de los distintos cálculos de los coeficientes mencionados en el párrafo anterior, en razón de que utiliza 2 técnicas: *a)* Siguiendo los procedimientos de la escala Lickert; y *b)* La técnica de la escala discriminatoria de Edwards y Kilpatrick. Vamos a desarrollar primero los procedimientos siguiendo las técnicas de Lickert y luego utilizaremos la técnica de Edwards y Kilpatrick.

A) La construcción de un escalograma Guttman

Los pasos a dar en la construcción de una escala de este tipo son: 1) Se construye una serie de ítems relevantes a la actitud que se quiere medir. 2) Se administran los ítems a una muestra de sujetos que van a actuar como jueces. 3) Se asignan puntajes a los ítems según la dirección positiva o negativa del ítem. 4) Análisis de ítems para la formación de series acumulativas. 5) En base a los ítems seleccionados se construye la escala final.

En los pasos 1, 2 y 3 se siguen los procedimientos señalados en la escala Lickert para la construcción de ítems, la administración a los jueces y la asignación de puntajes (ver escala Lickert).

4) El análisis de los ítems

Una vez aplicados los ítems a los jueces, procedemos al análisis de los ítems en su conjunto. La idea es formar una serie acumulativa de ítems. Para ello, vamos a diseñar un escalograma.

a) Primero computamos el puntaje total para cada uno de los jueces (es decir, sumamos los valores obtenidos en cada uno de los ítems).

b) Ordenamos a los jueces según el puntaje total, desde el puntaje más alto al puntaje más bajo.

Si nuestra serie de ítems fuese perfecta, todas las celdas cruzadas (*crossed cells*) en el escalograma estarían en una posición sobre una diagonal que corre desde el ángulo superior izquierdo hasta el ángulo inferior derecho. Mientras más alto sea el número de desviaciones a esta diagonal, más baja será la reproducibilidad y menos idéntica a una serie de ítems acumulativos será nuestra serie.

Antes de entrar en un análisis en detalle del cuadro 8, vamos a mostrar un escalograma más simple con el fin de clarificar la idea de *acumulación* y de *error*.

Supongamos que 20 jueces hayan respondido a 6 ítems en términos de “acuerdo” o “desacuerdo”. Ordenamos a los jueces, figurando en el cuadro 8 los “acuerdos” con una x, y los “desacuerdos” con una O. La distribución se muestra en el cuadro 8.

El cuadro 8 muestra una serie de ítems escalonados. El ítem 4 ocupa la primera posición en la escala en razón de que se dieron en él 3 respuestas positivas (de acuerdo);

el ítem 2 ocupa el segundo lugar ya que se dieron 6 respuestas positivas, y así sucesivamente hasta llegar al último ítem, el número 3, con el que acordaron 16 de los 20 jueces.

En términos de escalograma, las respuestas redondeadas con un círculo son “errores”, esto es, respuestas que caen fuera de la pauta general del escalograma señalado con la línea segmentada.

El cuadro analítico para el cuadro 8 sería:

Número de preguntas:	6
Número de jueces:	20
Cantidad total de respuestas:	$6 \times 20 = 120$
Cantidad de errores:	7

$$\text{Reproducibilidad} = 1 - \left(\frac{\text{Cantidad de errores}}{\text{Cantidad total de respuestas}} \right) = 0.94$$

Compliquemos un poco el ejemplo, utilizando 5 categorías o alternativas de cada ítem.

Supongamos que 30 jueces han juzgado 6 ítems en una escala de 5 puntos:

- 4— Muy de acuerdo
- 3— De acuerdo
- 2— Indeciso
- 1— En desacuerdo
- 0— Muy en desacuerdo

CUADRO 8. *Análisis de escalograma. Respuestas de 20 jueces a 6 ítems en términos de acuerdo-desacuerdo*

Rango del juez							Puntaje
1	X	X	X	X	X	X	6
2	X	X	X	X	X	X	6
3	X	X	X	X	X	X	6
4	0	X	X	X	X	X	5
5	0	X	⊙	X	X	X	4
6	0	0	X	X	X	X	4
7	0	0	X	X	⊙	X	3
8	0	0	X	X	X	⊙	3
9	0	0	0	X	X	X	3
10	0	0	0	X	X	X	3
11	0	0	0	X	X	X	3
12	0	0	⊗	0	X	X	3
13	0	⊗	0	0	X	X	3
14	0	0	0	0	X	X	2
15	0	0	0	0	0	X	1
16	0	0	0	0	0	X	1
17	0	0	0	0	⊗	⊙	1
18	0	0	0	0	0	X	1
19	0	0	0	0	0	0	0
20	0	0	0	0	0	0	0
	3	6	8	11	14	16	

Es decir, que teóricamente los puntajes van a variar desde un máximo de 24 puntos a un mínimo de 0.

El cuadro 9 señala la ordenación de los jueces, y el puntaje obtenido en cada uno de los ítems. Considerando la dispersión sobre todas las alternativas de la escala, vemos que la probabilidad de obtener una escala acumulativamente perfecta disminuye en la medida en que se trata ahora de que haya escalabilidad entre los ítems y en el interior de un ítem.

Si examinamos el ítem 1 encontramos que hay 8 personas con respuesta 4 (“muy de acuerdo”); 10 personas con respuesta 3 (“de acuerdo”); 4 personas con respuesta 2 (“indeciso”); 4 personas con respuesta 1 (“en desacuerdo”), y 4 personas con respuesta 0 (“muy en desacuerdo”). De haber escalabilidad perfecta en el ítem, deberíamos esperar que todos los jueces con respuesta 4 se coloquen en un orden de rango por encima de aquellos que dieron respuesta 3. No ocurre así, ya que sólo 3 de ellos están claramente por encima, mientras que hay sujetos con puntaje 3 e incluso un sujeto con puntaje 1 que aparecen en un orden de rango superior. Si observamos los sujetos con puntuación 2 vemos que están perfectamente ordenados en relación a los que tienen puntaje 3 y 4. En los sujetos con puntaje 1 vemos que existen dos sujetos mal ubicados (dos errores) en relación al ordenamiento escalable. De los 4 sujetos que tienen puntaje 0 en el ítem, 3 están fuera de posición.

CUADRO 9. Análisis de escalograma. Ordenación de puntajes y valores de respuesta en 6 ítems tipo escala Lickert (30 jueces)

Puntaje	Ítem 1					Ítem 2					Ítem 3					Ítem 4					Ítem 5					Ítem 6				
	4	3	2	1	0	4	3	2	1	0	4	3	2	1	0	4	3	2	1	0	4	3	2	1	0	4	3	2	1	0
24	X					X					X					X					X					X				
23	X					X					X					X					X					X	X			
21	X					X					X					X					X					X				
20		X				X					X					X					X					X				
19	X					X					X					X					X					X			X	
19	X					X					X					X					X					X				
19		X				X					X					X					X					X				
18	X					X					X					X					X					X				
18		X				X					X					X					X					X				
18		X				X					X					X					X					X				
18	X					X					X					X					X			X		X				
17		X				X					X					X					X					X				
16		X				X					X					X					X			X		X				
15		X				X					X	X				X					X					X				
15			X			X					X					X					X	X				X				
13		X				X					X					X					X			X		X				
12		X				X					X					X					X			X		X				
12		X				X					X					X	X				X			X		X				
11	X					X					X					X					X	X				X				X
10			X			X					X					X					X			X		X				
8				X		X					X					X					X				X	X				
8				X		X					X					X					X			X		X				
8				X		X					X					X					X	X			X					
8			X			X					X					X					X			X		X				
7			X			X					X					X					X			X		X				X
6				X		X					X					X					X			X		X				X
4			X			X					X					X					X			X		X				X
4				X		X					X					X					X			X		X				X
3				X		X					X					X					X			X		X				X
2				X		X					X					X					X			X		X				X

Podríamos hacer un escalograma con el ítem 1, de la misma manera como lo hicimos en el cuadro 1, y contar el número de errores.

Considerando la dispersión sobre todas las alternativas de la escala, la probabilidad de obtener una serie acumulativa de ítems en el ejemplo del cuadro 9 disminuye, en relación con el ejemplo del cuadro 8. Mediante la unificación de algunas categorías de respuesta de distintas maneras, podríamos tratar de aumentar la reproducibilidad de la escala.

Las técnicas que se utilizan para este propósito son diferentes. Se puede recurrir

directamente a máquinas electrónicas para procesamiento de datos, a alguna técnica de “tablas de escalograma” (*scalogram board techniques*), o un procedimiento más complicado y que implica pérdida de tiempo como el de “ensayo y error”.

Volviendo al ejemplo original del cuadro 9, tenemos errores o muchos sujetos fuera de posición dondequiera que pongamos la diagonal o los cortes. Se tratará entonces de combinar las respuestas de manera que podamos representar un nuevo *continuum* en el cual, por ejemplo, se combinen respuestas “muy de acuerdo” con respuestas “de acuerdo”, o cualquier combinación que indique escalabilidad. Esto significa que el criterio por el que nos guiamos para la combinación de respuestas es empírico, esperando que, si efectivamente el ítem discrimina, ordenará las pautas de discriminación con una combinación tal que minimice el error. Cuando el ítem no discrimina, es decir, cuando cualquier combinación produce un número grande de errores (y grande significa simplemente que el número de aciertos es menor que el número de errores, o que no existe una desproporción significativa entre aciertos y errores que nos dará como resultado un coeficiente de reproducción satisfactorio), el ítem no es escalable y se elimina.

En el ítem 1 podemos dejar las categorías 4 y 3 separadas, y combinar las categorías de respuesta 2, 1 y 0. En el ítem 2 combinamos las categorías de respuesta 4 y 3, por un lado, y 2, 1 y 0 por el otro, etcétera.

Supongamos que las combinaciones propuestas para cada ítem y sus respectivos pesos son ahora:

<i>Ítem</i>	<i>Combinaciones</i>			<i>Nuevos pesos</i>		
1	4	3	2,1,0	2	1	0
2	4,3		2,1,0	2		0
3	4,3		2,1,0	2		0
4	4,3,2		1,0	2		0
5	4,3		2,1,0	2		0
6	4,3,2		1,0	2		0

La distribución de los puntajes en los ítems adquiere ahora la siguiente forma:

CUADRO 10

Puntajes	Ítem 1			Ítem 2		Ítem 3		Ítem 4		Ítem 5		Ítem 6	
	2	1	0	2	0	2	0	2	0	2	0	2	0
12	X			X		X		X		X		X	
12	X			X		X		X		X		X	
12	X			X		X		X		X		X	
11		X		X		X		X		X		X	
11		X		X		X		X		X		X	
11		X		X		X		X		X		X	
11		X		X		X		X		X		X	
10	X			X		X		X			X	X	
10	X			X		X			X	X		X	
10	X			X		X		X	X	X		X	
10	X			X		X		X			X	X	
9		X		X			X	X		X	X	X	
9		X			X	X		X		X		X	
8			X	X		X			X	X		X	
7		X		X		X			X	X	X	X	
6	X				X		X		X	X	X	X	
5		X			X	X			X	X	X	X	
5		X			X	X	X		X	X	X	X	
5		X			X	X	X		X	X	X	X	
4			X		X	X	X	X		X	X	X	
4			X		X	X	X	X	X	X	X	X	
4			X		X	X	X	X	X	X	X	X	
2			X		X	X	X		X		X	X	
2			X		X	X	X	X			X	X	
2			X		X	X	X	X	X		X	X	
0			X		X	X	X		X	X	X		X
0			X		X	X	X		X	X	X		X
0			X		X	X	X		X	X	X		X
0			X		X	X	X		X	X	X		X
0			X		X	X	X		X	X	X		X

Obsérvese que, con la nueva clasificación de los ítems, el ítem 6 no contiene ahora ningún error, siendo a su vez el que menos discrimina entre los puntajes extremos. Los ítems 2 y 3 escalan bastante bien y con poco error. El ítem 4 tiene algunos errores, mientras que el 5, a pesar de la dicotomización, contiene más errores que los demás. La distribución de los puntajes en el ítem 1 indica la necesidad de unir las nuevas categorías 2 y 1 con el fin de reducir aún más el error. Tal como está el cuadro, su coeficiente de reproducibilidad será igual a:

CUADRO 11

Puntajes	Ítems				
	4	2	3	1	6
5	X	X	X	X	X
5	X	X	X	X	X
5	X	X	X	X	X
5	X	X	X	X	X
5	X	X	X	X	X
5	X	X	X	X	X
5	X	X	X	X	X
5	X	X	X	X	X
5	X	X	X	X	X
4	0	X	X	X	X
4	0	X	X	X	X
4	⊗	X	⊗	X	X
4	⊗	⊗	X	X	X
4	0	X	X	X	X
3	0	X	X	⊗	X
3	0	0	X	X	X
2	0	0	0	X	X
2	0	0	0	X	X
2	0	0	0	X	X
2	0	0	0	X	X
2	⊗	0	0	0	X
1	0	0	0	0	X
1	0	0	0	0	X
1	0	0	0	0	X
1	⊗	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0
0	0	0	0	0	0

$$r_p = 1 - \frac{\text{Total de errores}}{\text{Total de respuestas}} = 1 - \frac{38}{390} = .913$$

Valor satisfactorio, sin embargo, si unimos las categorías 1 y 2 del ítem 1, y eliminamos el ítem 5; detenemos entonces el escalograma de la página anterior, en el cual los ítems aparecen ordenados por grado de dificultad.

Puede verse ahora que los errores han sido reducidos en forma considerable, aunque a costa de reducir la discriminación en el interior de los ítems, además de eliminar el ítem originalmente designado con el número 5.

El lector notará que con los nuevos pesos asignados a los ítems se produce un reordenamiento en el rango de los jueces como consecuencia de la alteración en los puntajes totales.

Las líneas segmentadas representan los puntos de cortes.

B) La determinación de los errores

Hay dos métodos que tienden a dar resultados diferentes: la técnica de Cornell y la técnica de Goodenough.

Comencemos con la técnica de Cornell, la cual se realiza estableciendo puntos de separación en el orden de rango de los jueces, puntos que los separan en distintas categorías, tales como se definirían de ser la escala perfecta. Ésta es la técnica utilizada cuando señalábamos puntos de separación por medio de líneas segmentadas. Si observamos el cuadro 10, vemos que el ítem 1 necesita dos puntos de separación: el primero cae entre la última persona con puntaje 10 y la primera persona con puntaje 9. Todos los sujetos o jueces por encima de este punto de separación deberían tener puntaje 2 en el ítem. En términos de errores encontramos pues 5 errores (cuatro por encima y uno por debajo del punto de separación). El siguiente punto de separación lo establecemos entre la última persona con puntaje 5 y la primera persona con puntaje 4. No existe ningún error.

Finalmente, en el último punto de corte que corresponde al puntaje 0 tenemos 1 error. Así pues, el ítem 1 tiene un total de 6 errores. Puesto que hay 30 jueces y 3 categorías, el porcentaje de error es de 6.7%. Cuando el error es mayor que el 10% y el ítem permite reagrupación es necesario realizarla, o de lo contrario conviene eliminar el ítem. Se opera de la misma manera con los otros ítems.

Para el cálculo del coeficiente de reproducibilidad por la Técnica de Cornell, hay que seguir entonces los siguientes pasos: *a)* Determinar los *cutting points* de manera tal que se minimicen los errores. *b)* Ningún *cutting point* debe estar tan por encima o tan por debajo que la categoría menor incluya menos del 10% de los jueces. *c)* Ningún *cutting point* debe ser realizado de manera tal que el número de errores sea mayor que el número de jueces en la categoría.

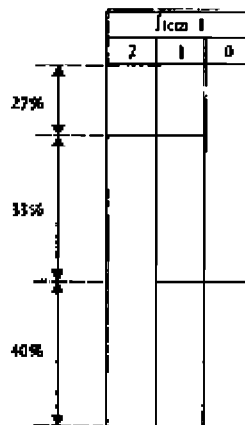
Para controlar si el lector entendió la manera de calcular el coeficiente de reproducibilidad por la técnica de Cornell, lo invitamos a hacer el cálculo para el cuadro 10 por su cuenta. Para los *cutting points* indicados en el cuadro, el coeficiente debe ser igual a .949. De hecho, habría que aclarar que el ítem 5 debería ser eliminado, con lo cual el coeficiente de reproducibilidad aumentaría a un valor aceptable, tal como vimos en el cuadro 11.

La técnica de Goodenough es un poco más complicada y se basa en el cálculo de errores en base a pautas marginales.

Los *cutting points* en esta técnica se determinan según la distribución de las respuestas de los jueces a cada una de las alternativas en los distintos ítems. Los cálculos para el cuadro 10 nos darían los siguientes valores:

	Ítems												
	1			2		3		4		5		6	
Categorías	2	1	0	2	0	2	0	2	0	2	0	2	0
Frecuencias	8	10	12	14	16	15	15	14	16	16	14	24	6
Porcentajes	27	33	40	47	53	50	50	47	53	47	53	80	20

Los *cutting points* para cada uno de los ítems deben seguir entonces los porcentajes respectivos. En el caso del ítem 1, los cortes serían entonces:



Es decir que, en el ítem 1, el primer corte o *cutting point* caería entre el primer sujeto con puntaje total 10 y el segundo con puntaje 10, y así sucesivamente para el resto de los ítems.

Mediante estos *cutting points*, se van a determinar pautas de respuestas correspondientes a cada corte. La pauta de respuesta es la manera “correcta” en que deberían distribuirse para cada puntaje total de cada juez si la escala tuviera perfecta escalabilidad. Cada respuesta que no sigue la pauta de respuesta ideal se considera un “error”. En nuestro ejemplo, un juez con puntaje de 8 debe seguir una pauta: 0-0-2-2-2-2. Un sujeto con puntaje 6 debe consecuentemente tener una pauta de respuesta 0-0-0-2-2-2. Es decir que primero hay que ordenar los ítems en términos de escalabilidad, para luego determinar las pautas correspondientes a cada valor. El principio de la técnica es bastante simple. Para 4 ítems dicotomizados y con valores de 1 y 0 para alternativas de respuesta “de acuerdo” y “en desacuerdo”, respectivamente, un puntaje total de 3 puede seguir cuatro tipos de pautas diferentes:

a) 0 1 1 1
b) 1 0 1 1
c) 1 1 0 1
d) 1 1 1 0

De estas pautas, solamente la primera es correcta, para ítems que están ordenados en forma acumulativa. Cada una de las pautas b), c) y d) tiene respectivamente 2 errores (uno por tener un 1 donde debería haber un 0, y otro por tener un cero donde debería tener un 1).

Ítems								Puntaje total	Pauta de respuesta	Errores			
1		2		3		4							
1	0	1	0	1	0	1	0						
X		X		X		X		4	1	1	1	1	0
	X	X		X		X		3					0
X			X	X		X		3					2
		X		X			X	3	0	1	1	1	2
	X	X		X		X		3					0
	X	X		X		X		3					0
	X	X		X			X	2					2
	X		X	X		X		2	0	0	1	1	0
	X		X	X		X		2					0
	X		X		X	X		1					0
	X		X		X	X		1	0	0	0	1	0
	X		X		X	X		1					0
	X		X		X		X	0					0
	X		X		X		X	0	0	0	0	0	0
	X		X		X		X	0					0
Total errores										6			

Para hacer más claro el método, en la página anterior dimos un ejemplo más simple, en el que presentamos 4 ítems dicotomizados, las pautas de respuesta y el cálculo de los errores.

El coeficiente de reproducibilidad se calcula según la fórmula ya conocida de:

$$r_p = 1 - \frac{\text{Cantidad de errores}}{\text{Número total de respuestas}}$$

En el ejemplo arriba citado:

$$r_p = 1 - \frac{6}{60} = .90$$

El coeficiente de reproducibilidad nos indica la proporción de respuestas a los ítems que pueden ser correctamente reproducidas.

C) Ejemplos

Veamos ahora dos ejemplos, a partir de los cuales presentaremos criterios adicionales al coeficiente de reproducibilidad para la determinación del universo de contenido en la escala Guttman.

Cuadro de la página siguiente: escalograma Guttman para la medición de nivel de vida determinada en función de materiales empleados en la construcción de viviendas y de las instalaciones sanitarias. Muestra de población de Colombres, en el departamento de Santa Cruz, en Tucumán, Argentina. Respuestas de 33 jefes de familia (arq. Hernández, CES; Universidad de Tucumán, Argentina, 1966).

REFERENCIAS

<i>Ejemplo B</i>	<i>Ejemplo A</i>
<i>Ítem 1:</i> Alejamiento de aguas servidas a pozo ciego y fosa séptica.	<i>Ítem 1:</i> Pileta de lavado.
<i>Ítem 2:</i> Agua corriente.	<i>Ítem 2:</i> Inodoro.
<i>Ítem 3:</i> Eliminador de residuos.	<i>Ítem 3:</i> Ducha.
<i>Ítem 4:</i> Techo de cinc, loza o teja con aislación.	<i>Ítem 4:</i> Revestimiento de cemento o superior en baño o letrina.
<i>Ítem 5:</i> Piso de cemento o mejor.	
<i>Ítem 6:</i> Luz eléctrica.	

<i>Ejemplo A: Vivienda</i>							<i>Ejemplo B: Instalaciones sanitarias</i>				
<i>Sujetos</i>	<i>Ítems</i>						<i>Sujetos</i>	<i>Ítems</i>			
	<i>1</i>	<i>2</i>	<i>3</i>	<i>4</i>	<i>5</i>	<i>6</i>		<i>1</i>	<i>2</i>	<i>3</i>	<i>4</i>
1	X	X	X	X	X	X	3	X	X	X	X
33	X	X	X	X	X	X	15	X	X	X	X
31	X	X	X	X	X	X	17	X	X	X	X
22	X	X	X	X	X	X	7	X	X	X	X
19	X	X	X	X	X	X	11	X	X	X	X
16	X	X	X	X	X	X	30	X	X	X	X
4	X	X	X	X	X	X	2	X	X	X	X
6	X	X	X		X	X	1	X	X	X	X
3		X	X	X	X	X	4	X	X	X	X
15		X	X	X	X	X	5	X	X	X	X
17		X	X	X	X	X	8		X	X	X
14		X	X	X	X	X	6		X	X	X
7		X	X	X	X	X	33		X	X	X
2		X	X	X	X	X	31		X	X	X
18		X	X	X	X	X	22		X	X	X
5		X	X	X		X	19		X	X	X
8		X	X		X	X	16		X	X	X
11		X	X	X		X	14		X	X	X
30		X	X	X		X	9				X
29			X		X	X	18				X
9				X	X	X	13				
13				X	X		25				
10				X	X		24				
26					X	X	29				
28					X	X	27				
24					X		26				
27						X	32				
20						X	12				
21						X	20				
25							28				
12							23				
23							10				
32							21				

Cálculo Universo de Contenido para el cuadro 11:

	<i>Ejemplo A</i>	<i>Ejemplo B</i>
Coefficiente de reproducibilidad:	.949	1.00
Rango marginal mínimo:	.72	.60
Alcance de distribución marginal:	.229	.40

El coeficiente de reproducibilidad en el ejemplo B, cuyo valor es 1.00, nos permite decir, sin tener en cuenta la distribución gráfica, que el sujeto 18 tiene en su casa únicamente revestimiento de cemento (o mejor, en su baño), mientras que el sujeto 33, cuyo puntaje es 3, tiene revestimiento, inodoro y ducha, pero no pileta de lavado.

Es muy difícil lograr escalabilidad perfecta, y consecuentemente existen errores que van a ser interpretados como errores de reproductividad. Guttman aconseja que los coeficientes de reproductividad no sean menores de .90.

El coeficiente de reproductividad (r_p) es un criterio necesario, pero no suficiente para la determinación de la escalabilidad de los ítems. Deben tomarse en cuenta otros factores. Stouffer *et al.*⁹ señalan cuatro criterios adicionales: 1) Alcance de la distribución marginal; 2) Pauta de errores; 3) Número de ítems en la escala, y 4) Número de categorías de respuestas.

1) Alcance de la distribución marginal

Es el más importante de los criterios adicionales, y debe acompañar al coeficiente de reproductividad. El criterio de distribución marginal es determinado por el Rango Marginal Mínimo (M. M. R.), que consiste en el r_p menos el promedio de los modos de las frecuencias relativas de las distribuciones de los ítems: ($r_p - MMR$).

Para algunos, los valores de este criterio adicional deben variar entre .15 y .35; para otros el mínimo debe ser mayor que .10. Estos valores indican *la escalabilidad de los ítems*, dato que no es proporcionado por el r_p de manera completa (es decir, es posible alcanzar valores altos de r_p —digamos .90— y resultar una escalabilidad inaceptable). Éste es el caso en el cual los *cuttings points* están muy próximos entre sí, con el resultado de discriminar solamente en los extremos de la escala y no a lo largo de la misma. En nuestro ejemplo, los valores de r_p son altos y muy aceptables; los alcances de la distribución marginal, en cambio, son aceptables para el ejemplo A, y demasiado altos para el ejemplo B.

2) Pauta de errores

Cuando el r_p es menor que .90, pero es escalable, es decir que tiene un r_p M. M. R. mayor que .10, estamos en presencia de más de una variable; mejor dicho, de una variable dominante y de una u otras menores, en el área a través de la cual se ordenan los sujetos. Este tipo de escalograma es denominado cuasi-escala. Éste no es el caso de los dos ejemplos que presentamos.

3) Número de ítems en la escala

A mayor número de ítems, mayor la seguridad de que el universo (del cual estos ítems son una muestra) es escalable. Es por esto que cuando los ítems están dicotomizados,

como es el caso en nuestros ejemplos, es aconsejable que su número sea mayor que 10. Pero puede usarse un número menor de ítems si las frecuencias marginales se colocan en un rango con recorridos del 30% al 70%.

En los ejemplos dados por nosotros el rango de frecuencias es:

<i>Ejemplo A</i>		<i>Ejemplo B</i>	
<i>Ítem 1</i>	24%	<i>Ítem 1</i>	30%
<i>Ítem 2</i>	57%	<i>Ítem 2</i>	55%
<i>Ítem 3</i>	60%	<i>Ítem 3</i>	55%
<i>Ítem 4</i>	69%	<i>Ítem 4</i>	60%
<i>Ítem 5</i>	78%		
<i>Ítem 6</i>	87%		

De acuerdo al requisito citado más arriba tenemos así alguna seguridad de que el universo se comporta como la muestra.

4) Número de categorías de respuestas

Es otro criterio para asegurar la escalabilidad; cuanto mayor el número de categorías, mayor la seguridad de que el universo es escalable. Por ello, a pesar de la necesidad de reducir las categorías por razones prácticas (disminución del número de errores), hay que asegurarse de que tal reducción no es la resultante de obtener frecuencias marginales extremas (.90-.10) que, como vimos más arriba, no permiten errores, pero artificialmente.

Si mantenemos el número de alternativas de respuestas, a pesar de que aumentará el número de errores, disminuimos la posibilidad de que aparezca una pauta escalable cuando de hecho el universo no lo es.

C) La técnica de la escala discriminatoria de Edwards y Kilpatrick

La selección de los ítems iniciales para una escala Guttman ha sido siempre un misterio. En su primer libro Guttman afirma que la selección es cuestión de intuición. Pero ¿cómo opera esta intuición? ¿No existen reglas o hechos que puedan ser usados cuando se seleccionan los ítems escalables? Aparentemente la intuición de Guttman opera satisfactoriamente porque muchas escalas que han usado esta “técnica” terminan por ser perfectamente escalables. Y aun en las publicaciones más tardías la vía de la intuición parece ser la más correcta.

Si nosotros recapitulamos el proceso total para el investigador, vemos que él, inicialmente, tiene un número de afirmaciones que más o menos miden adecuadamente la variable. De éstas seleccionará un pequeño número al cual aplicará el análisis de escalograma o tal vez la técnica H. Nada indica en la técnica de Guttman que los ítems deben ser elegidos de manera que representen diferentes pasos en el continuo psicológico. Como una cuestión de hecho, todos podrían tener una posición, digamos de

7 en un continuo tipo Thurstone. El único criterio es que los ítems sean escalables de acuerdo con los efectos discriminativos entre un grupo alto y un grupo bajo.

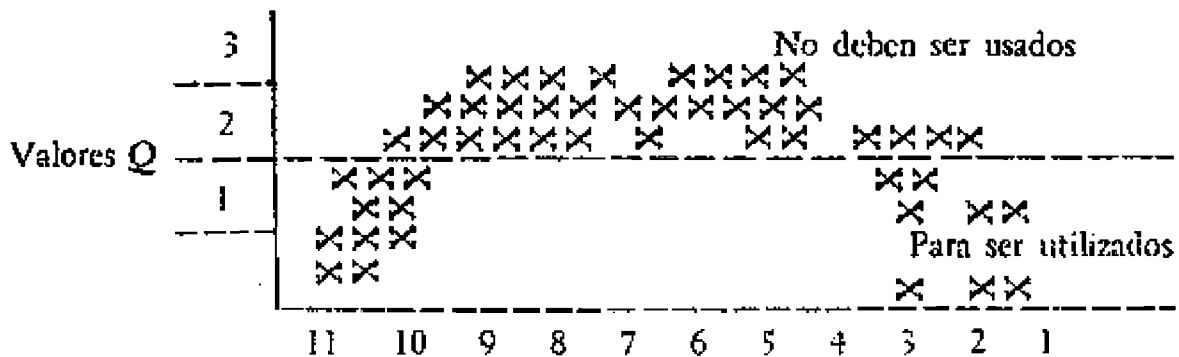
Para superar esto, se distribuyen representativamente los ítems a lo largo del continuo psicológico, y para establecer algunas reglas para la selección de los ítems iniciales de una escala Guttman, Edwards y Kilpatrick desarrollaron una técnica llamada: técnica de la escala discriminatoria (*the scale discrimination technique*).

Los principales pasos en el procedimiento podrían ser expuestos de la siguiente manera:

1) Por medio de la técnica de Thurstone de los intervalos aparentemente iguales (*the Thurstone equal appearing interval-technique*), un gran número de ítems son colocados a lo largo del continuo.

2) Se calcula para cada ítem la posición mediana y los valores Q.

3) Para disminuir el número de ítems se seleccionan aquellos que en los diferentes intervalos de la escala tienen valores bajos de Q; para esto deben colocar todos los ítems en una figura como la siguiente:



Se coloca el *cutting point* como indica la línea — (mediana de la escala Q) y se incluyen aquellos ítems que tienen valores bajos de Q; los que tienen valores altos se excluyen.

4) Ahora se vuelcan los ítems que quedan en una escala Likert, adicionando 5 o 6 alternativas de respuesta para cada uno.

5) Se aplica la nueva escala a por lo menos 100 jueces, y se separan el grupo alto (Q) y el grupo bajo (Q₁), de acuerdo con los puntajes totales.

6) Para probar el poder discriminativo de los ítems se dicotomizan los pesos como en el siguiente ejemplo:

	Grupo bajo	Grupo alto
(5)	2	4
(4)	8 <i>a</i>	3 <i>d</i>
(3)	7	1
(2)	28	32
(1)	32 <i>c</i>	38 <i>b</i>
(0)	19	17

La idea es seleccionar el *cutting point* en el continuo de los pesos con el fin de minimizar la suma entre las celdas *a* y *d*. Si colocamos el *cutting point* entre (3) y (4) nosotros tendríamos el siguiente resultado:

6	7
86	25

en donde la suma mínima sería $a + d = 31$. Si elegimos el *cutting point* entre (2) y (1) la suma mínima sería igual a 49 y, finalmente, si lo elegimos entre (3) y (2) encontraremos que la suma mínima es 25. Entonces seleccionamos esta alternativa como la mejor.

7) Después que se ha hecho la selección de las combinaciones mínimas para todos los ítems, se aplica un test r_ϕ a todos los ítems para seleccionar aquellos que tengan un alto poder discriminatorio. La fórmula a utilizar es la siguiente:

$$r_\phi = \frac{bc - ad}{\sqrt{(a + b)(b + d)(a + c)(c + d)}}$$

donde *a*, *b*, *c* y *d* se refieren a las siguientes celdas:

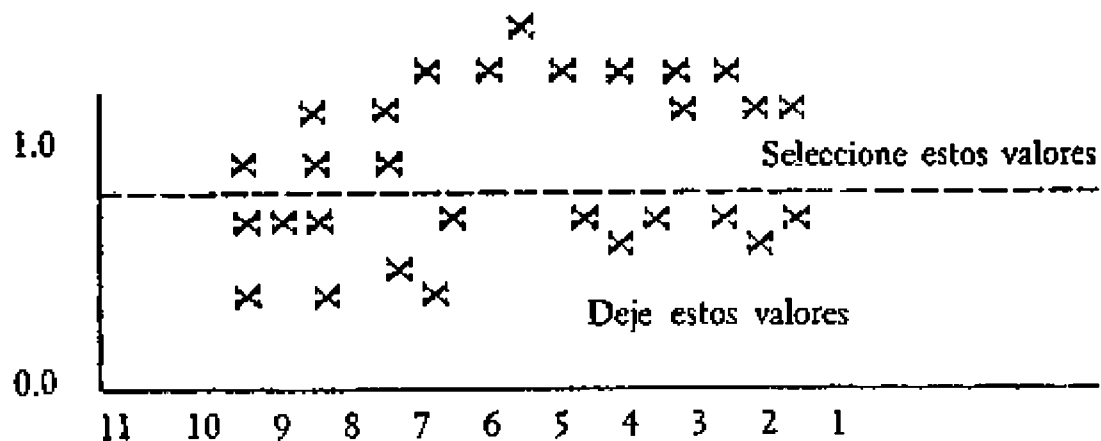
	Grupo bajo	Grupo alto	Total
	<i>a</i>	<i>b</i>	<i>a + b</i>
	<i>c</i>	<i>d</i>	<i>c + b</i>
Total	<i>a + c</i>	<i>b + d</i>	

Las figuras que aparecen más arriba reciben el siguiente valor r_ϕ

$$r_\phi = \frac{(87.79) - (17.8)}{\sqrt{(104)(95)(96)(87)}} = .73$$

que es un valor aceptable.

8) Ahora que se tienen los valores para todos los ítems se colocan todos los ítems en un diagrama para elegir los que tengan valores más altos de r_ϕ . El procedimiento es el mismo que se indicó para el punto 3.



Debe notarse que todo el tiempo tratamos de mantener el continuo psicológico subyacente intocable (manteniendo la distribución de los ítems sobre la escala de 1 a 11 más o menos homogénea).

CUADRO 12

Indiv.	Í T E M S										Score
	1	2	3	4	5	6	7	8	9	10	
1	1	1	4	1	1	3	3	3	4	4	25
2	3	4	3	3	3	4	4	4	4	4	36
3	4	4	4	1	2	2	2	4	4	4	31
4	2	3	1	2	3	4	4	3	3	4	29
5	1	3	2	3	2	4	4	4	4	4	31
6	3	3	3	3	2	4	4	4	4	4	34
7	3	4	3	1	3	3	4	3	4	4	32
8	1	3	1	3	2	4	4	4	4	4	30
9	3	4	3	1	2	4	3	4	4	4	32
10	4	4	4	4	4	4	4	4	4	4	40
11	0	0	0	1	0	3	3	3	0	4	14
12	4	4	4	3	3	4	4	3	4	4	37
13	1	3	1	1	0	2	0	4	4	4	20
14	4	4	4	4	4	4	4	4	4	4	40
15	3	1	1	0	0	2	1	0	3	2	13
16	1	1	1	1	0	2	2	3	2	0	13
17	0	2	2	1	1	3	3	4	4	4	28
18	3	4	4	2	4	4	3	4	4	4	36
19	3	3	4	0	0	4	4	4	4	4	30
20	2	4	3	3	3	4	4	4	4	4	35
21	2	4	3	3	3	3	4	4	4	4	34
22	3	3	4	3	3	4	4	4	4	4	35
23	3	3	4	3	3	4	3	2	2	2	29
24	1	4	4	3	1	4	3	4	4	4	32
25	3	4	3	2	4	4	4	4	0	4	32

D) La técnica H para mejorar la escalabilidad de la escala Guttman

Para mejorar la escalabilidad de la escala Guttman se puede utilizar la técnica H. La idea es formar por medio de esta técnica nuevos ítems (*contrived items*) agrupados con algunos de los ítems iniciales.

Supongamos que tenemos los datos que aparecen en el cuadro 12 (p. 192).

Es decir, 25 individuos han respondido a 10 ítems. Nosotros queremos aumentar la escalabilidad formando *contrived items*.

1) Existen cinco categorías de respuestas para cada ítem arbitrariamente numeradas 0, 1, 2, 3, 4.

2) Para cada persona se computa un puntaje total basado en todos los ítems que serán sometidos al análisis de escala.

3) Se obtienen las tablas de correlación de cada uno de los ítems considerados con el puntaje total provisorio. Para cada ítem se hacen distintas participaciones, considerando

como categorías positivas las respuestas números 1, 2, 3 y 4, después la 2, 3 y 4, después la 3 y 4 y luego la 4.

Empezamos por ejemplo con ítem 1, tomando 3 y 4 como positivo:

CUADRO 13

Puntaje total clasificado	0,1,2	Positivos 3,4
35 - 40	1	6
30 - 34	4	6
25 - 29	3	1
20 - 24	1	
15 - 19		
10 - 14	2	1
5 - 9		
0 - 4		

4) Seleccionamos los *cutting points* para cada ítem que se correlacionan con el puntaje (*score*) total, los cuales son lo suficientemente altos para formar una tabla de 2 por 2 en la cual ninguna celda error tiene una frecuencia mayor que la menor frecuencia de las dos celdas de la diagonal principal. Es decir, en este caso:

CUADRO 14

	0, 1, 2	llegamos a: 3, 4
35 - 40 30 - 34	a) 5	b) 12
25 - 29 20 - 24 15 - 19 10 - 14 5 - 9 0 - 4	c) 6	d) 2

→ Frecuencia Positiva = 12

Calculamos el r_Q :

$$r_Q = \frac{bc - ad}{\sqrt{(a+b)(b+d)(a+c)(c+d)}} = \frac{12 \cdot 6 - 5 \cdot 2}{\sqrt{17 \cdot 14 \cdot 11 \cdot 8}} = .43$$

Se siguen haciendo los mismos cálculos de r_Q para el ítem 1, pero ahora agrupando las categorías de otra manera. Las otras posibilidades son:

	<i>Positivas</i>
0	1, 2, 3, 4
0, 1	2, 3, 4
0, 1, 2	3, 4 (menos utilizado)
0, 1, 2, 3	4

5) Ordenamos todos los *cutting points* de la frecuencia positiva más alta a la más baja. Véase cuadro 15.

6) Seleccionamos conjuntos triples de ítems para constituir los cuatro nuevos *contrived* ítems. El objetivo principal es seleccionar ítems aceptables con la misma frecuencia aproximada para cada triple y espaciar cuanto sea posible estos conjuntos de triples tan extendidos (cuadro 15).

CUADRO 15

<i>Ítem</i>	<i>Categorías positivas</i>	<i>Frecuencia</i>	<i>Q</i>	
6	1234	25	.04	
6	234	25	.04	
7	1234	24	.08	
8	1234	24	.08	
8	234	24	.08	
10	1234	24	.08	
10	234	24	.08	
2	1234	24	.08	
3	1234	24	.08	
1	1234	23	.44	
4	1234	23	.44	
7	234	23	.44	
8	34	23	.44	
9	1234	23	.04	
9	234	23	.04	
10	34	22	.62	
10	4	22	.62	*
2	234	21	.73	*
6	34	21	.20	*
7	34	21	.73	} <i>Contrived Item I</i>
9	34	21	.60	
2	34	20	.80	
3	234	20	.85	
5	1234	20	.80	
9	4	19	.79	
3	34	17	.92	*
5	234	17	.88	*
1	234	16	.89	*
6	4	16	.77	} <i>Contrived Item II</i>
8	4	16	.53	
4	234	15	.81	
1	34	12	.43	
7	4	13	.63	
2	4	12	.97	*
4	34	12	.83	*
5	34	12	.94	*
3	4	10	.26	
5	4	5	.81	*
1	4	3	.53	*
4	4	2	.37	*
				} <i>Contrived Item IV</i>

Reglas: a) Si es posible, no repetir los ítems iniciales en diferentes *contrived items*. b) Buscar conjuntos de 3 ítems con un máximo de correlación (r_Q).

7) Adjudicar puntajes a cada individuo en cada *contrived item* asignándole un puntaje positivo si fue positivo en dos o tres de los ítems componentes del *contrived item*.

CUADRO 16. *Contrived items*

<i>Individuos</i>	<i>I</i>	<i>II</i>	<i>III</i>	<i>IV</i>
1	—			
2	—	—	—	+
3		+	+	+
4	—	—	+	+
5	—	—	—	+
6	—	—	+	+
7	—	+	+	+
8	—	—	—	+
9	—	—	+	+
10	+	+	+	+
11	—	—	—	+
12	—	+	+	+
13	—	—	—	+
14	+	+	+	+
15	—	—	—	—
16	—	—	—	—
17	—	—	—	+
18	—	+	+	+
19	—	—	+	+
20	—	+	+	+
21	—	+	+	+
22	—	+	+	+
23	—	+	+	+
24	—	+	—	+
25	—	+	+	+

* No escalable

CR = .96

MMR = .76

Entonces hemos llegado a 4 nuevos ítems (*contrived items*) que muestran un mayor coeficiente de reproductividad y una mayor diferencia entre CR y MMR.

E) La versión final de la escala

Hemos analizado el conjunto total de ítems, llegando finalmente a 10 ítems escalables según los criterios de la escala Guttman. Reproducimos aquí 4 de ellos y la manera de presentarlos en un cuestionario (en el ejemplo que citamos los 10 ítems han sido escalables para 5 valores en las alternativas de respuestas).

1. El patriotismo y la lealtad son los requisitos primeros y más importantes que debe llevar todo buen ciudadano. Ud. está:	☐ Totalmente de acuerdo. ☐ De acuerdo en general. ☐ Ni de acuerdo ni en desacuerdo. ☐ En desacuerdo en general. ☐ Totalmente en desacuerdo.
2. Quien no quiere pelear por su país merece algo mucho peor que la cárcel o los trabajos forzados. Ud. está:	☐ Totalmente de acuerdo. ☐ De acuerdo en general. ☐ Ni de acuerdo ni en desacuerdo. ☐ En desacuerdo en general. ☐ Totalmente en desacuerdo.
3. Cualquier esfuerzo hecho en el entrenamiento militar de Chile es compensado por los beneficios de su seguridad. Ud. está:	☐ Totalmente de acuerdo. ☐ De acuerdo en general. ☐ Ni de acuerdo ni en desacuerdo. ☐ En desacuerdo en general. ☐ Totalmente en desacuerdo.
4. Para Chile sería un grave error permitir el ingreso de los extranjeros que quitan las oportunidades de trabajo a los nacionales. Ud. está:	☐ Totalmente de acuerdo. ☐ De acuerdo en general. ☐ Ni de acuerdo ni en desacuerdo. ☐ En desacuerdo en general. ☐ Totalmente en desacuerdo.

La escala es unidimensional, es decir que los resultados obtenidos de los sujetos van a permitir ubicarlos en el continuo favorable-desfavorable a la actitud, en un rango que va (en nuestro ejemplo) de 0 a 40 puntos.

F) Ventajas y desventajas de la escala Guttman

Ventajas. a) Se asegura en forma casi definitiva la unidimensionabilidad de la escala. b) Conociendo el puntaje de un individuo se puede saber qué grado de acuerdo tuvo con los ítems y ubicarlo así en el continuo de la escala.

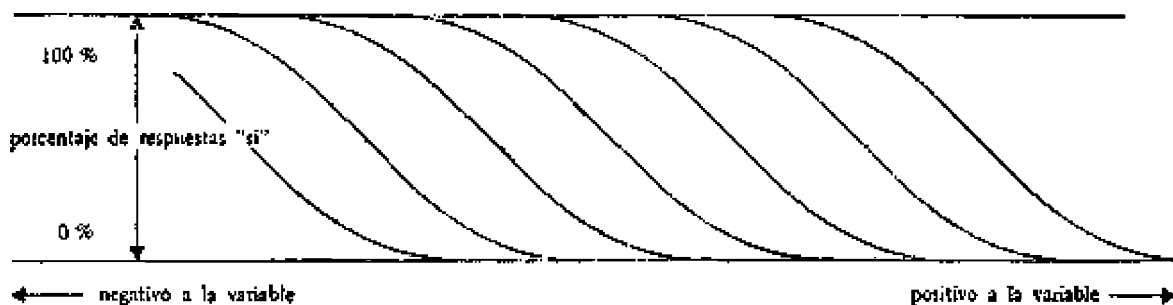
Desventajas. a) Cuando se trata de medir actitudes complejas conviene hacer un tipo

de escala para cada dimensión de la actitud. *b)* Involucra mayor cantidad de trabajo que en las dos escalas mencionadas anteriormente (Lickert y Thurstone). *c)* Escalas de este tipo pueden resultar unidimensionales para un grupo y no para otros grupos.

G) Comentarios finales

Los ítems de la escala Guttman, como los de la escala Lickert, son acumulativos. Sin embargo, los ítems en la escala Lickert son, cada uno de ellos, acumulativos; mientras que en la escala Guttman consiste en una *serie acumulativa* de ítems, donde cada ítem tiene además carácter de acumulativo. De ahí que en la escala Lickert sea posible que todos los ítems puedan ocupar aproximadamente el mismo lugar de la escala; en el caso de la escala Guttman, utilizando la técnica H de Edwards y Kilpatrick, garantizamos que los ítems en la escala se distribuyan a lo largo del continuo de actitud. En términos gráficos:

FIGURA 8



Ejemplo de algunos ítems acumulativos que forman una serie acumulativa de ítems. Todos los ítems van en la misma dirección (todos los ítems son de tal tipo que todos los sujetos con una disposición negativa respondan "Sí"). En realidad la inclinación para cada ítem y las distancias entre los ítems varían.

¿Variación en el instrumento, en los sujetos, o en ambos?

La escala Guttman ordena tanto los sujetos como los estímulos con respecto a un continuo de actitud, es decir que tanto a sujetos como a estímulos se les puede asignar valores de escala. Consecuentemente, la escala Guttman corresponde a un enfoque centrado en la respuesta. En este caso, tanto la actitud del sujeto como la actitud reflejada por el estímulo actúan para determinar la respuesta del individuo.

MÉTODO DE COMPARACIÓN POR PARES

Consiste en presentar los estímulos al sujeto de a dos por vez (pares) y preguntarle cuál de ellos es el más grande, el mejor, el más cálido, etc. Por medio de este artificio logramos un orden de rango basado en el número de elecciones recibidas por cada ítem.

Como las comparaciones son sistemáticas (el sujeto compara todo estímulo con cada

uno de los otros), el resultado final nos indicará cuán consecuente es el individuo con sus juicios o evaluaciones. En otras palabras, esto nos da la consistencia de las respuestas o la confiabilidad. En una pauta totalmente consistente, debemos esperar que el estímulo que sea más grande, más agradable, más cálido, etc., para el sujeto tendrá $n - 1$ elecciones, el segundo estímulo $n - 2$ elecciones, y así sucesivamente hasta llegar al menos grande, que tendrá 0 elecciones.

Se debe tener cuidado con no presentar sistemáticamente un estímulo en primer o segundo lugar, ya que con ello podríamos favorecer o contrariar la elección. Para evitar esto, se colocan los estímulos aleatoriamente, o se administran formas paralelas.

Este método de comparación por pares requiere mucho tiempo cuando las comparaciones a realizar son numerosas. El número posible de pares a calcular resulta de la fórmula

$$\frac{n(n-1)}{2}$$

donde n = cantidad de ítems.

Es decir que si queremos comparar 50 ítems, el número de combinaciones que el sujeto deberá evaluar son: 1225; si utilizamos 100 ítems, el número de pares a comparar serán 4950.

A) Ejemplo de construcción de una escala basada en el método de comparación por pares

El método de comparación por pares puede ser usado como un test de Tipo A, o como un test de Tipo B. Analizaremos aquí los resultados de un test de Tipo B en el caso de los jueces y un test de Tipo A en la aplicación a un sujeto.

1) Versión de los jueces. Test de Tipo B

Supongamos que queremos determinar el grado relativo de izquierdismo-derechismo de 5 partidos políticos en un país cualquiera, para después poder ordenar el grado de izquierdismo-derechismo en sujetos según sus simpatías políticas por partidos.

En la versión de los jueces vamos a ubicar a los partidos en un orden de rangos, y además —según la evaluación de los jueces— en una escala intervalar en donde un extremo de la escala representará izquierdismo y el otro derechosismo.

A un grupo de 50 jueces solicitamos que ordenen los 5 partidos *A, B, C, D* y *E* dando un valor 0 al partido más a la izquierda, un valor 1 al siguiente, etc., hasta un valor 4 al partido ubicado más a la derecha del espectro político.

Analizaremos los resultados de 5 jueces (se entiende que la evaluación de la escala debe hacerse sobre el total de jueces que se utilicen) para simplificar las operaciones. Los resultados de los ordenamientos de nuestros 5 jueces son los siguientes:

<i>Juez</i>	<i>Ordenación de los partidos</i>				
	A	C	B	D	E
1	4	3	2	1	0
2	3	2	4	0	1
3	2	4	3	1	0
4	4	2	3	1	0
5	3	4	2	1	0

El puntaje máximo posible con 5 jueces para un partido es $\frac{20}{5} = 4$. El puntaje mínimo es 0.

Los puntajes para los cinco partidos son entonces:

$$A = (4 + 3 + 2 + 4 + 3)/5 = 3.2$$

$$C = 15/5 = 3.0$$

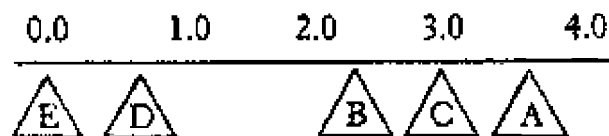
$$B = 14/5 = 2.8$$

$$D = 4/5 = 0.8$$

$$E = 1/5 = 0.2$$

El orden de rango de los partidos, de derecha a izquierda, es:

$A - C - B - D - E$, y su ubicación en una escala intervalar:



2) Versión de la escala para sujetos. Versión final. Test de Tipo A

En el cuestionario se presentan los partidos políticos por pares, en una lista en la que se combinen de dos en dos todos los partidos. En nuestro caso los pares a comparar serían:

A-B	B-D
A-C	B-E
A-D	C-D
A-E	C-E
B-C	D-E

Cuidamos, por supuesto, que los partidos aparezcan “mezclados” en la lista, evitando por ejemplo que el partido A esté siempre al comienzo, y se le solicita al sujeto que seleccione entre cada par a cuál de los dos partidos prefiere.

B) Ejemplo de una escala de comparaciones por pares en un cuestionario (entrevista)

El entrevistador muestra una tarjetita (hoja III) separada con el contenido indicado con la línea — — — — arriba. En el cuestionario aparece el texto total arriba. El entrevistador sigue las instrucciones preguntando cada vez por un solo par.

Las indicaciones se hacen naturalmente en el espacio en el cuestionario y no en la tarjeta (hoja III). (El entrevistador hace el registro en el cuestionario señalando el entrevistado solamente su preferencia.)

Un entrevistado ha contestado de la siguiente manera:

35. Tengo aquí una hoja con los nombres de los partidos políticos, agrupados de dos en dos. Por favor, dígame cuál partido prefiere en cada grupo <i>Muestra hoja III</i>	
Subraye el partido escogido, si dos partidos son escogidos al mismo tiempo, subraye ambos. Si el entrevistado no puede escoger ningún partido, no subraye nada	<u>Conservador</u> - <u>Liberal</u> <u>Socialista</u> - <u>Comunista</u> <u>Liberal</u> - Partido Laborista <u>Laborista</u> - <u>Socialista</u> <u>Laborista</u> - <u>Conservador</u> <u>Comunista</u> - <u>Laborista</u> <u>Conservador</u> - <u>Socialista</u> <u>Liberal</u> - <u>Comunista</u> <u>Socialista</u> - <u>Liberal</u> <u>Comunista</u> - <u>Conservador</u> Conserv. 1 = 1 Lib. 2 + 2 + 2 + 2 = 8 Cent. 2 + 2 + 1 = 5 Soc. 2 + 2 = 4 Com. 2 = 2

Las ponderaciones son calculadas después por el entrevistador. Se pone 2 si una de las alternativas fue tomada, y 1 si ambas fueron tomadas. El resultado es el perfil 1, 8, 6, 4, 1 del sujeto. Podemos concluir que la persona simpatiza más con el partido liberal. La ponderación 2 y 1 que corresponden a cada par es arbitraria. Se puede utilizar, por ejemplo, también 1 y 0. En este caso el perfil sería 0, 4, 3, 2, 0, y el sujeto todavía simpatiza más con los liberales.

Además de poder determinar la posición del sujeto en el *continuum* izquierda-derecha, se puede, por ejemplo, estudiar las personas que tienen *inconsistencia* en sus perfiles. Además del perfil obtenido podemos calcular la posición del individuo según nuestra escala:

Resultados para el individuo

<i>Posición del partido a lo largo de la variable</i>	<i>Nombre del partido</i>	<i>Columna M Núm. de elecciones "recibidas"</i>	<i>Posición del individuo a lo largo de la variable</i>
++ 4	A	1	4 por 1 = 4
+ 3	C	8	3 por 8 = 24
2	B	5	2 por 5 = 10
- 1	D	4	1 por 4 = 4
-- 0	E	2	0 por 2 = 0
			Total = 42

El individuo tiene una posición medianamente en el centro de la escala. Una persona extrema en la variable recibe como máximo 60 puntos y como mínimo 20. Compruebe esto usted mismo a través de cálculos.

Es evidente que resulta mucho más fácil para el sujeto indicar cuál de las 2 alternativas prefiere más. Sopesar cada ítem en una escala es más difícil, como también es complicada la tarea de ordenar por rango los ítems. Cuando el número de ítems es más de 3 o 4 en la tarea de ordenar por rango, el sujeto mismo tiene que usar alguna especie de comparación de los ítems entre sí. De este modo el método de comparaciones por pares se hace en una forma sistemática. El resultado de este método es un orden de rango basado en el número de elecciones recibidas por cada ítem.

Una ventaja con el método de comparaciones por pares es que muestra cuán "consecuente" es el sujeto en su estimación de las alternativas comparadas. Una inconsecuencia se muestra al disminuir irregularmente la serie de números en la columna *M*. Si el número de ítems es N y si el sujeto es consecuente, nosotros obtendremos $(N - 1)$ elecciones recibidas para el ítem que al sujeto le gusta más, $(N - 2)$ para su segunda preferencia de ítem, etc. Si el número de ítems es 7, así recibimos 6, 5, 4, 3, 2, 1, 0 (en el caso de la persona mencionada anteriormente). 6, 5, 4, 2, 2, 2, es un ejemplo de inconsecuencia. La suma de los valores de la serie de números es en ambos casos 21, el número de comparaciones por pares.

Hay algunas variantes del método de comparaciones por pares, como el método de comparaciones por pares dobles y comparaciones tríadas que han sido usadas en medición de intereses y actitudes. Estos métodos, sin embargo, hacen más difícil que el sujeto haga una ordenación "consecuente" (precisión más baja).

Algunos *test* de interés y actitudes muy conocidos que usan el método de orden por rango o el método de comparaciones por pares son: el *test* de interés de Allport Vernon, Lindzey en "*Study of Values*" (actitudes hacia la religión, política, arte, ciencia, etc.) y el registro de preferencia Kuder (interés en actividades como deporte, música, etcétera).

C) Ventajas y desventajas

Ventajas. a) Proporciona resultados más precisos b) Muestra claramente cuán “consecuente” es un sujeto con sus propios juicios. c) Es más fácil para un sujeto indicar cuál de dos alternativas prefiere más.

Desventajas. a) Cuando los juicios son numerosos, las comparaciones son muy grandes. En el caso de 20 ítems, el número de comparaciones a realizar es de 190. b) Gran consumo de tiempo de entrevista.

EL DIFERENCIAL SEMÁNTICO. LA ESCALA DE OSGOOD

A) Antecedentes teóricos

El método es descrito por los autores como un método para medir el significado que tiene un objeto para un individuo.

Osgood supone que existe un espacio semántico de dimensionalidad desconocida y de naturaleza geométrica. El espacio está construido (o constituido) de escalas semánticas. Cada escala consista de un par de adjetivos que son bipolares. Se supone que estas escalas forman una función lineal que pasa a través del origen. Para estar en condiciones de definir el espacio adecuadamente, es necesario usar una gran cantidad de escalas que son una muestra representativa extraída del universo de escalas. Para diferenciar el significado de un objeto, el individuo hace una elección entre las alternativas dadas. La función de cada elección es localizar el objeto en el espacio semántico. La validez de la localización en este punto en el espacio depende del número y representatividad de las escalas.

De este modo, la diferencia semántica significa la estabilización sucesiva (anclaje) de un objeto hasta un punto en el espacio multidimensional semántico, a través del puntaje de un número de alternativas semánticas dadas presentadas en la forma de escalas. Una diferencia de significado entre 2 objetos es simplemente una función de las diferencias de su ubicación en el mismo espacio, es decir, una función de la distancia multidimensional entre 2 puntos.

El punto en el espacio que da una definición operacional del significado tiene 2 características principales: 1) Dirección desde el origen; 2) Distancia desde el origen. Esto podría ser explicado como el *tipo* e *intensidad* del significado.

La *dirección* desde el origen depende de cuál de los polos de la escala se elige y la *distancia* depende de cuán extrema es la posición elegida en la escala.

B) Dimensiones en el espacio semántico

Osgood dio gran importancia al muestreo. El diferencial semántico está influido por 3 fuentes de variación: el individuo, las escalas y los objetos. Muchas diferentes modificaciones fueron hechas para asegurar la universalidad de la estructura del factor (*factor-structure*) pero siempre Osgood obtuvo los mismos factores principales en los diferentes análisis y así llegó a la conclusión de que la estructura del factor no dependía de la elección de escalas. El seguir 3 factores de hecho explicó la mayor parte de la

varianza total, mientras otras dimensiones sólo explicaban una pequeña parte de ellas.

Dimensiones

1) La *evaluación* que hace el individuo del objeto o concepto que se está clasificando. Ejemplo de escalas bipolares: regular-irregular; limpio-sucio; bueno-malo; valioso-despreciable.

2) La percepción del individuo de la *potencia* o poder del objeto o concepto. Escalas: grande-chico; fuerte-débil; pesado-liviano.

3) La percepción del individuo de la *actividad* del objeto o concepto. Escalas: activo-pasivo; rápido-lento; frío-caliente.

C) Construcción

El método para el diferencial semántico no es una prueba con ciertos ítems y puntajes de *tests* específicos. Debe ser visto como un método para reunir cierto tipo de información (un método que puede ser generalizado), el cual tiene que constituirse por las demandas que presenta cierto problema de investigación. No hay objetos estándar o escalas estándar.

Selección del objeto (concepto): “objeto” se usa para determinar qué significa el “estímulo” que da “reacción” (respuesta) en el individuo a través de su indicación en las escalas de adjetivos.

El objeto puede ser verbal; puede constar de sólo una palabra o de varias palabras. Objetos no-verbales pueden ser diferenciados (cuadros u otros estímulos estéticos).

La elección correcta de un objeto es un problema de muestreo. Esto generalmente significa en la práctica común que el investigador usa su sentido común al seleccionar el objeto. El investigador debería pensar en elegir objetos: 1) Que se supone darán diferencias individuales (para poder estudiar la variación en el material). 2) Que tengan un solo significado (de otra manera hay riesgo de vacilación en la elección). 3) Que se supone que todos los individuos conocen bien (de otro modo habrá regresión al medio de la escala).

D) Selección de escalas

Al seleccionar las escalas se debe estar seguro de tener los 3 factores: evaluación, potencia y actividad representadas. Para cada factor uno puede seleccionar cierta cantidad de escalas correlacionadas y a través de éstas obtener el promedio de las escalas. Esto aumentará la representatividad y la confiabilidad. Los promedios son llamados *factor scores*. El criterio más importante para la selección de escalas así es su patrón factorial (generalmente se seleccionan 3 escalas para cada factor. Las escalas elegidas cargadas en ese factor son bajas en los otros).

Otro criterio es la relevancia de las escalas al objeto. Las escalas con composición factorial desconocida pueden ser usadas si se suponen muy relevantes para la

investigación. En este caso uno puede incluir escalas con composición factorial conocida y usarlas como *escalas de referencia* y estudiar cómo las escalas “desconocidas” covarían con las escalas de referencia.

Análisis de los datos

Se deberán asignar pesos a las diferentes alternativas de la escala: 1, 2, 3, 4, 5, 6, 7; o 7, 6, 5, 4, 3, 2, 1, dependiendo de cuál polo en el par de adjetivos debe ser considerado como alto para este factor.

Por ejemplo:

_____	_____	_____	_____	
bueno	malo	o	bueno	malo

dependiendo del orden de presentación. Las escalas dentro del mismo factor se presentan como variable positiva y negativa como en la escala Lickert. Usando los pesos es posible hacer cálculos aritméticos corrientes, perfiles, análisis factorial, etcétera.

E) El test de Osgood y la medición de actitudes

Tal como se ha hecho notar, el *test* de Osgood no es una escala de actitudes en el sentido corriente, sino que las investigaciones han demostrado que los valores de escala pueden ser usados para la medición de actitudes. Se mide la actitud general pero no el contenido de la actitud en el significado denotativo.

Osgood piensa que con la ayuda de los valores de escala es posible formar una escala universal que podría ser usada para medir la actitud hacia cualquier objeto. La ventaja con tal escala universal sería: 1) Económica (tiempo, dinero, trabajo). 2) Disponible y así hacer posible la medición de actitudes que no habían sido previstas. 3) Posibilidades de comparar entre diferentes investigaciones de actitud y actitud hacia diferentes objetos.

A continuación presentamos un *ejemplo*:

INSTRUCCIONES

Le presentamos algunos párrafos tomados de periódicos y revistas que indican una relación entre diferentes asuntos de interés. Su tarea consistirá en evaluar dichos párrafos de acuerdo a una serie de escalas de adjetivos que encontrará a continuación del texto que *usted* debe juzgar.

El procedimiento consiste en señalar con una X en la escala el lugar en el cual usted ubicaría *su* opinión acerca del texto.

Por ejemplo: La palabra *fanático* podría ser evaluada en las siguientes escalas:

flexible	___	:	___	:	___	:	___	:	___	:	___	:	___	rígido
negativo	___	:	___	:	___	:	___	:	___	:	___	:	___	positivo
delicado	___	:	___	:	___	:	___	:	___	:	___	:	___	rudo

Si alguien considerara que la palabra *fanático* indica rigidez, debería marcar en la escala el lugar que está más próximo al adjetivo ríido, de este modo:

flexible	___	:	___	:	___	:	___	:	___	:	___	:	<u>X</u>	rígido
-----------------	-----	---	-----	---	-----	---	-----	---	-----	---	-----	---	----------	---------------

Si considerara que indica rigidez pero en menor grado, marcaría en:

flexible	___	:	___	:	___	:	___	:	___	:	<u>X</u>	:	___	rígido
-----------------	-----	---	-----	---	-----	---	-----	---	-----	---	----------	---	-----	---------------

o en cualquier otro lugar menos cercano al adjetivo ríido, según lo considerara. Lo mismo haría con cada una de las otras escalas. *Haga usted lo mismo con cada uno de los textos.*

Recuerde que debe señalar *su* opinión. No deje de hacerlo en ninguna de las escalas.

Si tiene preguntas fórmúelas antes de comenzar a trabajar. Es conveniente que trabaje en silencio y lo más rápidamente posible.

Porque obviamente se ve que una América Latina desarmada espiritual y materialmente ante la amenaza creciente y reciente de la Conferencia Tricontinental quedará atada de pies y manos para la invasión, la subversión y el consiguiente desembarco de infantería de marina en la lucha anticomunista posterior, inmediata e inevitable.

bueno	___	:	___	:	___	:	___	:	___	:	___	malo
débil	___	:	___	:	___	:	___	:	___	:	___	fuerte
exaltado	___	:	___	:	___	:	___	:	___	:	___	moderado
absurdo	___	:	___	:	___	:	___	:	___	:	___	significativo
blando	___	:	___	:	___	:	___	:	___	:	___	duro
pasivo	___	:	___	:	___	:	___	:	___	:	___	activo
positivo	___	:	___	:	___	:	___	:	___	:	___	negativo
poderoso	___	:	___	:	___	:	___	:	___	:	___	impotente
frío	___	:	___	:	___	:	___	:	___	:	___	cálido

La mísera situación social de la minoría negra, que vive hacinada en casas de madera de los barrios extremos, ha rebelado una vez más a la gente de color que constantemente

lucha por mejorar sus condiciones de vida.

Frente a los aparatos de refrigeración, el aire acondicionado, la vivienda sana y cómoda de los blancos, los negros sólo disponen de las bocas de riego para refrescarse del espantoso calor de los últimos días. La policía no atendió razones y cerró las bocas de riego, pero los negros enfurecidos atacaron a los intrusos y esto dio origen a una más de las tantas tragedias raciales que se producen en los Estados Unidos.

bueno	___:	___:	___:	___:	___:	___:	___	malo
débil	___:	___:	___:	___:	___:	___:	___	fuerte
exaltado	___:	___:	___:	___:	___:	___:	___	moderado
absurdo	___:	___:	___:	___:	___:	___:	___	significativo
blando	___:	___:	___:	___:	___:	___:	___	duro
pasivo	___:	___:	___:	___:	___:	___:	___	activo
positivo	___:	___:	___:	___:	___:	___:	___	negativo
poderoso	___:	___:	___:	___:	___:	___:	___	impotente
frío	___:	___:	___:	___:	___:	___:	___	cálido

El mercado bursátil

Poca actividad y precios irregulares mostró durante las operaciones de ayer el mercado bursátil. Los mejores conformes pertenecieron a Banco de Chile, Tierra del Fuego, Copec, Melón, Papeles y Cartones, Cap. y Mademsa.

Los bonos se anotaron en pequeñas partidas de Deuda Interna 7-1 y Banco Hipotecario de Chile 10-6 con cotizaciones mantenidas.

El volumen de los negocios ascendió a la suma de E° 331 nominales en bonos y 279.582 acciones.

El movimiento en transacciones se vio poco activo siendo inferior al de la semana próxima pasada.

A la hora del cierre de las operaciones, la plaza mostraba poca actividad y precios irregulares.

bueno	___:	___:	___:	___:	___:	___:	___	malo
débil	___:	___:	___:	___:	___:	___:	___	fuerte
exaltado	___:	___:	___:	___:	___:	___:	___	moderado
absurdo	___:	___:	___:	___:	___:	___:	___	significativo
blando	___:	___:	___:	___:	___:	___:	___	duro
pasivo	___:	___:	___:	___:	___:	___:	___	activo
positivo	___:	___:	___:	___:	___:	___:	___	negativo
poderoso	___:	___:	___:	___:	___:	___:	___	impotente
frío	___:	___:	___:	___:	___:	___:	___	cálido

ESCALA DE DISTANCIA SOCIAL DE BOGARDUS

Para finalizar nuestro capítulo sobre escalas, presentaremos la escala de distancia social de Bogardus, quien de hecho fue el primer científico social en utilizar el concepto de escala para la medición de actitudes. La escala de distancia social fue utilizada por Bogardus para el estudio de las relaciones —manifestadas en términos de actitudes— entre grupos étnicos y nacionales. Ulteriormente la utilización de este tipo de escala se ha extendido al estudio de relaciones entre clases sociales, grupos ocupacionales, etcétera.

El concepto de distancia social se refiere al grado de conocimiento o simpatía o intimidad que caracterizarían las relaciones entre distintos grupos. Por medio de ella es posible medir el grado de “aceptabilidad” que un grupo social cualquiera tiene respecto de otro grupo. Los ítems a partir de los cuales se construye la escala están graduados de manera tal que van de la aceptación (medida por un grado alto de aceptación de relaciones íntimas) al rechazo (medido por la negación a cualquier relación).

A) Procedimientos básicos para la construcción de una escala de distancia social

1) Defina el continuo —dentro del sentido del concepto de distancia social arriba mencionado— en términos que van desde un contacto estrecho, afectuoso e íntimo, pasando por la indiferencia, hasta llegar a la antipatía, la hostilidad y el rechazo.

2) Reúna una serie de ítems de distancia social igualmente espaciados. Tales ítems constituyen afirmaciones que se refieren a un gran número de relaciones sociales.

3) De acuerdo con el método de Thurstone —unas 100 personas que hacen de jueces—, las juzgan en términos de grados crecientes de distancia social o de grados crecientes de proximidad social.

4) En el estudio de Bogardus, siete de estos ítems igualmente espaciados, de grados crecientes de distancia social según la determinación de los jueces, fueron seleccionados.

5) A estos siete ítems se les dan números arbitrarios que van de 1 a 7, respectivamente, representando creciente distancia social en esta proporción.

Debemos aclarar que Bogardus, para la selección de los ítems, muchas veces ha empleado su experiencia; es decir que no ha recurrido a jueces para que, de acuerdo a las evaluaciones de éstos, sean seleccionados finalmente los ítems considerados adecuados para que integren la escala.

En la práctica, lo que han hecho los investigadores es aplicar directamente la escala ya estructurada por Bogardus, consistente en 7 ítems. O, por el contrario, solamente han adaptado tales ítems —en sus formulaciones— a las necesidades del estudio.

*B) A continuación presentamos los ítems que ilustran esta escala
de distancia social*

ESCALA DE DISTANCIA RACIAL

Instrucciones

1) Dé su primera reacción ante cada pregunta sin pensarlo demasiado.

2) Dé sus reacciones a cada raza considerada como un grupo. No dé sus reacciones de sentimientos a los mejores o a los peores miembros que usted ha conocido, sino piense de acuerdo a la idea que usted tiene de la raza considerada como un todo.

3) Ponga una cruz debajo de cada raza en tantas de las siete filas como sus sentimientos le dicten.

Categoría (ítem)	Judío	Negro	Indio	Blanco europeo	Chino	Mestizo
1. Se casaría con						
2. Tendría como amigos regulares						
3. Trabajaría en una oficina junto a						
4. Viviría en el mismo vecindario con						
5. Estaría simplemente hablando como a conocidos a						
6. Excluiría de mi vecindario a						
7. Excluiría de mi país a						

El sujeto es, pues, interrogado para dar sus primeras reacciones de sentimientos y no para racionalizar. La suposición es que las primeras reacciones de sentimientos reflejan

actitudes mejor que ningún otro, aparte de la conducta en sí en un periodo.

Aunque la conducta de larga duración es la mejor prueba de las actitudes de una persona, la escala de distancia social se ideó para dar predicciones, mientras se espera que la conducta de larga duración revele las actitudes. La conducta de un corto periodo puede revelar pseudoactitudes, no actitudes verdaderas. Puede medir actitudes que son “ocultas” para ciertos fines.

Se supone, en consecuencia, que las primeras reacciones de sentimientos sin racionalización son significativas en revelar cómo actuaría una persona si tuviera que enfrentarse repentinamente con las situaciones citadas en la escala.

C) Flexibilidad de la técnica

Otra indicación de la flexibilidad de esta técnica puede verse en el hecho de que puede ser usada no solamente para escalar grupos o valores que son externos al sujeto que hace el *rating*, sino también para escalar a los mismos *raters* con respecto a la distancia social entre ellos y algún grupo.

Para ello se utiliza un *continuum* en el sentido de favorabilidad-desfavorabilidad de la escala tipo Lickert. Con un ejemplo se ilustrará mejor:

INSTRUCCIONES

1) Hablando de los norteamericanos, por favor, ¿podría indicarnos su simpatía o no hacia ellos?

2) Por favor, evalúe a los norteamericanos en esta escala, marcando con una cruz en los espacios punteados la afirmación que expresa su sentimiento hacia ellos.

No les tengo
ni simpatía
ni antipatía

Me son deci-
didamente
simpáticos

Me son un
poco simpá-
ticos

Decididamente
me son anti-
páticos

Les tengo un
poco de anti-
patía

.....

Esta forma de escala es recomendable y se presenta primero a los sujetos a quienes se les está sometiendo a investigación acerca de sus actitudes. Luego viene la forma anteriormente diseñada de la escala.

La forma ilustrada anteriormente se construye de la misma manera para los demás grupos étnicos, o nacionales, hacia quienes los sujetos tienen actitudes que nos interesa conocer.

D) Confiabilidad

Una escala de distancia social no es fácilmente probada para la determinación de su confiabilidad, ya sea por la forma múltiple o por la técnica de la división por mitades (*split-half*). El enfoque test-retest es la medida más efectiva de confiabilidad de tal

escala.

E) Validez

Para mostrar la validez de la escala se requiere pensamiento cuidadoso. La aplicación del método del *known-group* implicaría hallar grupos conocidos que sean favorables con algunos de los tipos étnicos y no favorables con otros. Si las respuestas de estos grupos forman el requisito patrón, entonces la validez parecería probable.

Por otro lado, el uso del método de criterios independientes requeriría que el orden de rango formara algún otro rango de aceptabilidad social. Tales indicadores podrían ser el orden de rango de deseabilidad como inquilinos en un gran proyecto de construcciones de viviendas, de aceptabilidad como miembros de un gremio, etcétera.

F) Limitaciones y aplicaciones

Esta técnica de *scaling* no está limitada en cuanto a su flexibilidad de aplicación ni por su crudeza como medida. Los principales problemas son: 1) La suposición de la equidistancia entre los puntos de la escala. 2) La suposición de que cada punto está necesariamente “más allá” del punto anterior. 3) El hecho de que puede ser probada por confiabilidad solamente por el “decir” de la técnica ordinaria del test-retest. Por lo tanto, el uso de este método de *scaling* está limitado generalmente a estudios pilotos o a investigaciones que por alguna razón deben ser completados rápidamente y no requieren un nivel de precisión muy alto.

BIBLIOGRAFÍA RECOMENDADA PARA ESCALAS DE MEDICIÓN DE ACTITUDES

Al lector interesado en profundizar tanto en la teoría como en la técnica de escalas de medición de actitudes recomendamos especialmente los siguientes textos:

Edwards, A., *Techniques of Attitude Scale Construction*; Appleton-Century-Crofts, Nueva York, 1957.

Torgerson, W., *Theory and Methods of Scaling*; J. Wiley, Nueva York, 1958.

Thurstone, L., *The Measurement of Attitudes*; The University of Chicago Press, Chicago, 1929.

Stouffer, S., *et al. Studies in Social Psychology in World War II. Measurement and Prediction*; Princeton University Press, 1950.

Upshaw, H., “Attitudes Measurement”, en Blalock, H., y Blalock, A.: *Methodology in Social Research*; McGraw-Hill, Nueva York, 1968.

Selltiz, Jahoda, *et al., Research Methods in Social Relations*; Holt-Rinehart-Winston, Nueva York, 1959.

Recomendamos también revisar especialmente las siguientes revistas:

Sociology and Social Research
Journal of Applied Psychology
Journal of Social Psychology
Journal of Abnormal and Social Psychology
American Sociological Review
American Journal of Sociology
Psychological Bulletin

VII. EL TRABAJO DE CAMPO

INGVAR AHMAN

SUPONGAMOS que los entrevistadores han sido seleccionados y se ha verificado su capacidad para el trabajo en el terreno en el cual va a ser realizado. Supongamos también que el investigador cuenta con una central para el trabajo de terreno, desde la cual él opera. Esta central tiene servicio telefónico, una gran sala de conferencias y algunas salas pequeñas, a las cuales los entrevistadores pueden llevar a los R ,¹ si es que existen dificultades para obtener un lugar adecuado para realizar la entrevista en otra parte. En la central existen, además, planos de la ciudad, guías telefónicas, mapas de comunicaciones, etcétera.

Algunos días antes de que comience el trabajo en el terreno, los entrevistadores deben tener una reunión preliminar, con el objeto de tomar todas las instrucciones necesarias, ejercitarse en algunos cuestionarios, formularios, etc. Algunos de los puntos de importancia que deben tratarse en estas reuniones preliminares son los siguientes:

1) Descripción del estudio

El investigador ofrecerá una descripción metodológica del estudio y explicará algunas partes teóricas que hay dentro de él.

2) Descripción de la muestra

El investigador dará una explicación acerca de la forma como se ha diseñado la muestra, su tipo y qué posibilidades existen para remplazar los individuos que no se encuentren al efectuar las entrevistas.

Los puntos 1 y 2 son meramente informativos para ubicar a los entrevistadores en el tipo de estudio, y si el investigador lo considera conveniente, pueden eliminarse. Los puntos que vienen a continuación son bien específicos y conviene tomarlos en cuenta.

3) Cómo ponerse en contacto con los respondentes

A cada entrevistador debe asignársele un distrito especial, propio. El entrevistador recibe un mapa y se le pide que haga un plan para sus contactos, que no ofrezca mucha pérdida de tiempo en movilización. De acuerdo con la profesión del R , el entrevistador selecciona el tiempo que considera más probable para contactarse con las personas elegidas. Es importante que se contacte con el mayor número de R posibles. No siempre existe la seguridad de que la persona disponga del tiempo necesario para la entrevista

cuando es visitado por el entrevistador, aunque esté realmente interesado. En este caso, el entrevistador deberá fijar un día y hora para la entrevista. Cuando establezca la hora para la entrevista, debe considerar el tiempo necesario para movilizarse, de tal modo que una entrevista no tenga que hacerse en forma apresurada, debido a que el tiempo para la próxima entrevista está demasiado cerca de la anterior. El entrevistador debe tener especial cuidado de no dejar nunca de asistir a una entrevista que ha sido previamente concertada. Si por algún motivo surgen dificultades imprevistas, el entrevistador debe acudir de inmediato a la central de operaciones a pedir cooperación. Cualquier otro entrevistador puede hacer la entrevista. En caso negativo, es necesario que el entrevistador llame a la persona y trate de concertar una nueva entrevista. Si el *R* no estuviere en casa al efectuar el contacto, el entrevistador puede dejarle un mensaje (esto podría estandarizarse para todos los entrevistadores), manifestando que el entrevistador se ha hecho presente y que ha estado buscando al *R* para fijar una nueva entrevista, en una fecha determinada, etc. En todo caso, es necesario que el entrevistador trate de evitar el contacto con el *R* a través del teléfono. Es mucho más fácil rehusar una entrevista por medio del teléfono que si el entrevistador lo solicita personalmente. De igual modo, no es recomendable contactar al *R* en su trabajo, ya que podría estar muy ocupado y rehusar la entrevista por esa causa.

La presentación

El entrevistador deberá presentarse a sí mismo al *R*, diciéndole, por ejemplo: “¿Cómo está usted? Mi nombre es... Estoy tratando de efectuar unas entrevistas para un estudio que hará la universidad. Este estudio desea conocer la opinión pública en asuntos de actualidad muy importantes. Usted ha sido seleccionado en la muestra para efectuarle una entrevista y me agradaría hacerle algunas preguntas, como ésta:...”

Hay ocasiones en que el *R* se pone suspicaz y desea saber algo más acerca del estudio. El entrevistador puede ofrecerle más detalles sobre el estudio *en general*, y su importancia para el progreso de la comunidad, etc. Si el respondente está aún titubeante, hay que partir con las preguntas: “¿Por qué no comenzamos con algunas preguntas y así usted se da una idea mejor?”

A veces el *R* no entiende por qué el entrevistador no escoge otra persona para ser entrevistada. Puede que sugiera que en vez de él se entreviste a un hermano suyo, que, por ejemplo, está más interesado en eso. Llama a su hermano que está en la habitación contigua y que está deseoso de ser entrevistado. La situación se torna difícil. El entrevistador tiene que explicar que las personas fueron escogidas previamente al azar y que no pueden ser sustituidas. Ese *R* y nadie más es la persona requerida. En muchos casos es muy importante explicar el objetivo del estudio y que sus resultados serán publicados y remitidos después al entrevistado. Es importante insistir en que las respuestas a las preguntas serán usadas solamente en tablas estadísticas y que la entrevista es estrictamente confidencial. Algunas de estas explicaciones conviene mencionarlas al final de la entrevista, de acuerdo a la forma como reacciona el *R*.

En la presentación, es evidente que el entrevistador debe mostrar sus credenciales que lo autorizan para efectuarla. Esto debe hacerse como una regla general, sean o no solicitadas por el *R*. La identificación personal, naturalmente, debe llevarse consigo. El entrevistador puede preguntar también al *R* si ha recibido alguna carta previa explicándole el estudio que se va a realizar. Si no la hubiere recibido, puede mostrarle una copia extra al *R*.

He aquí un ejemplo de una carta previa enviada con motivo de una entrevista:

INSTITUTO DE INVESTIGACIONES EN EL CAMPO SOCIAL
UNIVERSIDAD DE

Junio de 1975

Estimado señor (a):

El Departamento de Ciencias Sociales de la Universidad de está llevando a cabo una encuesta de la opinión pública en asuntos actuales de importancia.

El instituto auspicia muchos tipos de estudio. La presente encuesta es parte del estudio de la participación de la comunidad y de los ciudadanos, en toda la zona norte del país.

En las próximas semanas estaremos entrevistando tanto a hombres como mujeres en la zona norte. Su nombre ha sido seleccionado como una de las personas que serán entrevistadas. Las personas han sido elegidas al azar.

El resultado de todas las entrevistas será publicado en un informe que presentará al Norte como un todo. El informe será totalmente estadístico y ninguna persona será identificada. Todas las entrevistas serán tratadas en forma estrictamente confidencial.

Con el objeto de que la muestra sea exacta, no podemos efectuar sustituciones de las personas que han sido seleccionadas para las entrevistas.

Un entrevistador con sus credenciales y correcta identificación, que lo acredita como uno de los entrevistadores para este estudio, lo visitará en su casa en las próximas semanas. Estamos seguros de que usted encontrará esta visita interesante y útil.

Si usted tiene algunas preguntas que formularnos o desea información adicional acerca de nuestros estudios, nos será muy grato contestar a sus requerimientos.

Le agradecemos de antemano su gentil atención.

Atentamente,
Director del Estudio

Cuando se trate de la presentación, deben explicarse además al entrevistador los aspectos de la entrevista que dicen relación con la *cooperación o negativa* de parte del *R*. ¿Cómo puede el entrevistador ganarse la confianza del entrevistado para realizar la tarea en la cual está empeñado? ¿Podemos garantizarle esta confidencialidad? Precisamente para esto se utilizará este estudio. ¿Por qué se ha elegido exactamente al *R*? ¿Cómo haremos frente a los diferentes tipos de preguntas que el *R* frecuentemente hace al comienzo de la entrevista o antes?

Confidencialidad

Los entrevistadores han de ser entrenados debidamente acerca de la ética de la entrevista. El entrevistador no debe discutir ninguna entrevista con nadie que no pertenezca al equipo de investigación. Debe tomarse especial cuidado con el objeto de que los entrevistadores no discutan acerca de las entrevistas entre ellos en lugares públicos o donde puedan ser escuchados por personas que no pertenecen al equipo de investigación.

Los entrevistadores no deben mostrar el cuestionario ni lleno ni vacío a ningún otro *R* ni mucho menos a otras personas. El entrevistador firmará un recibo por el número de cuestionarios que ha recibido y ese número de cuestionarios deberá coincidir con el número de entrevistas que aparecen en la central, después que el trabajo haya terminado. Los cuestionarios sin llenar no pueden darse a ninguna persona que no pertenezca al equipo de investigación, salvo que cuenten con el visto bueno del investigador jefe.

¿Qué podemos hacer para obtener un alto grado de ética en el equipo de entrevistadores?

Pagos elevados, tanto de entrevistas efectuadas como también de aquellas que no han resultado, pero en las cuales el entrevistador ha hecho contacto con el *R*.

Uso de grupos especiales. Si el investigador usa un grupo de estudiantes interesados en el área de estudio, la motivación para un buen trabajo será alta. En algunos países existen equipos profesionales de entrevistadores que poseen título profesional y hay una ética de grupo debidamente establecida. Cualquier intento de hacer un mal trabajo, en este caso, significa graves perjuicios para el entrevistador.

Comprobaciones. El investigador efectuará comprobaciones al azar a los entrevistadores. Esto puede hacerse en la forma de un llamado telefónico al *R*, para saber cómo encontró la entrevista y hacerle diversas preguntas: “¿Cómo era el entrevistador? ¿Adecuado o falto de interés?” El investigador puede incluso efectuar visitas y repetir algunas preguntas a algunos de los *R* y preguntarles acerca de cómo encontraron al entrevistador.

En ambos casos, el investigador preguntará cuánto demoró la entrevista. En el caso en que algunas personas no hayan seguido las reglas indicadas, se les excluye de inmediato del equipo.

Informes a la central de entrevistadores

El procedimiento para los informes a la central debe ser fijado de tal modo que el investigador pueda verificar cómo se está efectuando el trabajo en el terreno. La primera entrevista en el terreno o las dos primeras deben ser verificadas inmediatamente por el investigador. Si hay algunos malos entendidos (incluso aunque el entrevistador haya sido entrenado antes en el cuestionario), deben corregirse inmediatamente. Puede ser que algunas preguntas en el cuestionario no estén suficientemente claras o que las instrucciones de los entrevistadores no sean claras o que el entrevistador no haya interpretado correctamente algo. Hay entonces una posibilidad de corregir esto inmediatamente sin destruir el estudio. Después de esto, por ejemplo, cada 5 entrevistas el entrevistador debe entregarlas al investigador para verificarlas.

El examen del contenido debe hacerse cuidadosamente. Si el entrevistador se ha saltado algunas preguntas, debe ir nuevamente con el *R* y obtener los datos que faltan. En algunos casos es necesario recordar a los entrevistadores algunas instrucciones y en la mayoría de los casos debe recordárseles escribir en una forma legible.

Instrucciones para la investigación

En cada trabajo de terreno debe dársele a los entrevistadores un folleto de instrucciones especiales acerca del estudio. El contenido de este folleto puede ser el siguiente:

- a) Información general sobre el estudio.
- b) Procedimientos de muestreo utilizados.
- c) Datos sobre la entrevista: *i)* Plan de obtención de datos para cada entrevistador; *ii)* Material con que debe contar; *iii)* Contacto con las autoridades locales; *iv)* Entrevistas para la práctica; *v)* Presentación a los *R*; *vi)* Plan de contactos con los *R*; *vii)* El contacto con la central; *viii)* Cómo motivar al entrevistado; *ix)* Confidencialidad del estudio.
- d) Códigos generales utilizados: *i)* Instrucciones para las preguntas.

Especialmente la parte *d* debe ser seguida por todos los entrevistadores en las reuniones. Si se usa algún código especial como el que hemos sugerido con *O*, *X* y *Y*, deberá explicarse en detalle y practicarse. Después de esto, todas las preguntas del cuestionario deben ser seguidas con las instrucciones especiales que se han dado para ello.

A continuación, los entrevistadores tomarán varios días para estudiar todo el material y volver antes de que comience el estudio para una comprobación de sus conocimientos.

Verificación en el cuestionario

El entrevistador deberá ser supervisado con el objeto de comprobar que conoce el cuestionario sin ninguna dificultad, y que conoce la instrucción existente para todas las preguntas y otras instrucciones generales. En el momento de la entrevista es demasiado tarde para empezar a buscar las instrucciones, ya que ésta debe tener cierta fluidez.

En general, una encuesta significa papeles y más papeles. Cuando el entrevistador sale al terreno debe llevar, por ejemplo, un equipo como el siguiente: 5 cuestionarios; 1 folleto con instrucciones para los entrevistadores; 5 juegos de tarjetas, de 3 tarjetas c/uno; 20 tarjetas para *R* (para identificación del *R*); 1 plano; credenciales; 3 cartas de presentación; 20 hojas especiales para dejar a los *R* que no se encuentren en la casa.

La inscripción de los R

En esta parte del capítulo nos proponemos dar algunas indicaciones de un sistema que controle la inscripción o registro de los *R* antes, durante y después del trabajo en el terreno.

Supongamos que el tamaño de la muestra es de 500 y hemos seleccionado las 500 personas al azar. Registramos por cada persona: nombre, profesión, edad, estado civil,

dirección y número de teléfono, la parte del país en que habita y su nacionalidad. Todos estos datos se anotan en una página estándar (por cada individuo se emplea una página). En esta página registramos, además, dónde se consideró a este individuo en la muestra. (Por ejemplo: provincia del Norte, Libro de Registro, pág. 345, ind. Núm. 56.)

Supongamos que usamos 25 entrevistadores. Por cada entrevistador anotamos una página normal donde se colocan sus *R*. En este caso, 20 por entrevistador, seleccionados de tal manera que vivan bastante cerca uno de otro. En esta hoja dejamos el espacio suficiente para registrar qué entrevista ha efectuado el entrevistador, cuáles han faltado y por qué y si las entrevistas se han entregado después a otros entrevistadores.

Ahora, para cada individuo hacemos un juego de dos tarjetas (exactamente iguales), de tamaño pequeño, de cartulina. En esta tarjeta colocamos el nombre del *R*, profesión, edad, estado civil, nacionalidad, dirección y número telefónico. Además, colocamos el nombre y número del entrevistador.

Esto es, cada entrevistador irá al terreno con 20 tarjetas para *R* y una copia de cada una se guardará en la central de entrevistadores. Llamemos a la tarjeta de los entrevistadores *Te* y a la de la central, *Tc*.

En la tarjeta *Te*, el entrevistador hace una anotación cuando la entrevista está lista y la razón, en caso de haber sido imposible efectuarla.

Tendremos dos tarjeteros para las inscripciones que hemos de manejar. Los llamaremos *A* y *B*. En *A* colocaremos las entrevistas terminadas y los *R* que no vamos a volver a molestar. En *B* guardaremos todas las *Tc* al comienzo y todas las *Te*, con las cuales tenemos que continuar. Usaremos los registros siguientes:

Al comienzo, cuando partimos al terreno encontramos todas las *Tc* en el tarjetero *B* bajo el entrevistador al que pertenecen (los números 1-6, etc., son los números de los entrevistadores).

Supongamos que un entrevistador vuelve con una entrevista que ha efectuado. Entonces la tarjeta *Tc* en *B* se retira de su lugar numerado y se pone con la tarjeta *Te* bajo las letras en *A*. Las letras en *A* se refieren a la inicial del apellido del *R*, en orden alfabético.

Ahora, supongamos que al entrevistador le ha sido imposible efectuar la entrevista durante el curso de tres semanas, porque el *R* se encuentra fuera de la ciudad. En el caso en que proyectemos emplear otro entrevistador en la última oportunidad, tomaremos las tarjetas *Te* y *Tc* y las colocaremos en el tarjetero *B* bajo el título "Fuera de la ciudad". Lo mismo podría hacerse en el caso en que alguien esté en el hospital por largo tiempo, haya concertado con el entrevistador una nueva entrevista para después o si alguien ha rehusado cooperar en la entrevista. En el caso en que la persona haya rehusado, el investigador puede tratar de probar otra vez con otro entrevistador. Esto es, todas las tarjetas que se encuentran en el espacio superior en el tarjetero *B* y las tarjetas que vale la pena intentar de nuevo.

Tan pronto como una entrevista esté lista, dijimos que su tarjeta pasa al tarjetero *A*. Lo mismo se hace en el caso en que se haya desistido de hacer la entrevista. Por ejemplo, la persona que no desea cooperar y que ha sido imposible convencerla; algunas personas

que se encuentran fuera de la ciudad o fuera del país hasta fines de año; otras que se encuentran en el hospital por un tiempo largo, es preferible no seguir con el caso. Algunas puede que sea imposible entrevistarlas por otras razones. Puede tratarse de un extranjero que vive en el país sin conocer el idioma y que haya sido imposible conseguir un intérprete; alguien que puede haber fallecido; personas que se encuentran en prisión, etcétera.

Hay muchas ventajas en el uso de este sistema que a primera vista puede parecer complicado. He aquí algunas de ellas:

En cualquier momento del estudio, el investigador tiene la situación completa del trabajo en terreno para una *rápida* inspección. Puede darse cuenta de inmediato de cuántas entrevistas hay hechas y cuántas están todavía por hacerse, con el simple control de los tamaños de las tarjetas en la parte de abajo de los tarjeteros *A* y *B*. Sabe inmediatamente cuántas entrevistas ha efectuado cada entrevistador, con sólo verificar el número del entrevistador en el tarjetero *B*. Puede cambiar a nuevos entrevistadores inmediatamente con sólo distribuir las tarjetas que se encuentran en la parte de arriba del tarjetero *B*. En este caso, las tarjetas *Te* van al nuevo entrevistador y las *Tc* quedan anotadas en el número del nuevo entrevistador. El investigador puede asimismo verificar cuántas entrevistas fue imposible efectuar con sólo mirar la parte superior del tarjetero *A* y compararla con las tarjetas que figuran en la parte baja del mismo tarjetero. De ello obtiene, también, el porcentaje de pérdida.

Aún más útil es el sistema si pensamos cuán fácilmente los “*R* están perdidos” entre diferentes entrevistadores, cuando se usan otros sistemas. Nadie sabe quién tiene un determinado *R*, en el momento en que se necesita conocer el dato, así como otras inconsistencias. Con el sistema presentado aquí tenemos asimismo una doble comparación de las tarjetas. Esto es, todas las tarjetas que van en el tarjetero *A* o *B* (parte superior) tienen que estar pareadas ($Tc + Te$). Esto es, ningún entrevistador puede andar con algunas tarjetas que están terminadas y ninguna tarjeta *Tc* puede existir cuando ya están terminadas las entrevistas.

En seguida, el investigador completa las “páginas *R*” y las “páginas del entrevistador” (mencionadas al comienzo), en base a los resultados de los dos tarjeteros. Incluso, si la mayor parte del trabajo en el terreno ha sido hecho inmediatamente, habrá un gran lapso en el cual las últimas entrevistas han sido efectuadas. En este caso el sistema presentado es también muy aconsejable y se adecua de inmediato para cualquier operación que se requiera.

VIII. ANÁLISIS DE DATOS: EL CONCEPTO DE PROPIEDAD-ESPACIO Y LA UTILIZACIÓN DE RAZONES, TASAS, PROPORCIONES Y PORCENTAJES

CARLOS BORSOTTI

DESPUÉS de haber considerado los problemas lógicos (conceptualización, formulación de hipótesis, clasificación, variables, etc.) y lógico-empíricos (operacionalización de variables, niveles de medición, índices, etc.), así como las distintas técnicas de recolección de datos y los problemas involucrados en ellas (entrevista, cuestionario, codificación, etc.), es el momento de considerar dos grandes unidades temáticas: *A)* Cómo disponer los datos para decidir si la hipótesis puede ser aceptada o debe ser rechazada. *B)* Cómo manejar los datos matemáticamente para aceptar o rechazar la hipótesis.¹

A) EL CONCEPTO DE PROPIEDAD-ESPACIO

El concepto de propiedad-espacio intenta dar una respuesta al primer tema.

Como es sabido, en sociología es difícil aplicar rigurosamente un diseño experimental, constituyendo grupos experimentales y de control estrictamente controlados. Pero eso no significa que la lógica subyacente al diseño experimental no sea utilizable.

Siempre que se disponen los datos para decidir acerca de si una hipótesis puede ser aceptada o debe ser rechazada, se está siguiendo la lógica de un diseño experimental. Es decir, se constituyen grupos que poseen distintos grados de una dimensión (educación baja, educación alta; ingreso bajo, ingreso medio, ingreso alto) y se somete a todos los grupos a un mismo estímulo (una pregunta o un conjunto de preguntas de un cuestionario). Se espera que los distintos grupos constituidos respondan de manera distinta al mismo estímulo. Si no responden de manera distinta, significa que los distintos grados de la dimensión sobre cuya base se constituyeron los grupos no tienen ninguna relevancia con relación al estímulo al que fueron sometidos los grupos. Sobre esto se volverá a insistir, ya que es un aspecto fundamental.

Toda hipótesis relaciona variables propiamente tales entre sí, o variables atributos entre sí. Para mayor comodidad llamaremos variables tanto a las variables propiamente dichas como a las variables atributos, ya que la diferencia parece centrarse, principalmente, en los niveles de medición de ambos tipos.

Como los niveles de medición son cuatro (nominal, ordinal, interval, de razones o proporciones), las posibles combinaciones de niveles de medición tomando dos variables

son 10 (nominal-nominal, nominal-ordinal, nominal-interval, nominal-de razones, ordinal-ordinal, ordinal-interval, etc.). A partir de esto, debe tenerse presente que el nivel de medición resultante de la combinación de dos o más variables siempre es el nivel de medición adecuado a la variable de más bajo nivel de medición. Es decir, si se combinan una variable nominal con una ordinal y una de razones o proporciones, las operaciones estadístico-matemáticas permitidas no pueden ser otras que las permitidas para el nivel de medición nominal.

Todo el tema de la propiedad-espacio está dominado por estos tres temas aludidos: hipótesis, diseño experimental y nivel de medición. Existen ciertas consideraciones subyacentes sobre probabilidades y estructura latente, a las que en este momento no se va a hacer referencia pero que conviene tener presentes.

Se puede caracterizar la propiedad-espacio como el conjunto de las posiciones que pueden tener las unidades en estudio en una variable o en un sistema de dos o más variables. Una variable puede tener distinto número de valores posibles. Cada uno de ellos es una posición disponible para las unidades en estudio. El conjunto de estas posiciones disponibles es la propiedad-espacio en la que se ubicarán las unidades. Es conveniente aclarar que las unidades en estudio pueden ser individuos, colectivos de individuos y colectivos de colectivos.

La propiedad-espacio en una variable

Una variable define un espacio de una sola dimensión y, por lo tanto, puede ser representada por una línea, aunque esto no es del todo cierto cuando se analizan las variables según su nivel de medición.

Si la variable es de un nivel de medición nominal o variable atributo, no corresponde, en realidad, hablar de una línea, sino de un conjunto de compartimientos separados, estancos, alineados, sin que la alineación signifique ninguna clase de supuesto jerarquizado entre los compartimientos. Si, por ejemplo, se toma la variable nacionalidad, puede suponerse un conjunto de cajas alineadas sin contacto entre sí, en cada una de las cuales se ubican las unidades en estudio.

Si la variable es de un nivel de medición ordinal, se tiene nuevamente un conjunto de compartimientos estancos, separados, alineados. Pero aquí el hecho de ubicar una unidad en uno de los compartimientos significa acordarle un valor menor, igual o mayor que otras unidades. Por eso es que la alineación de los compartimientos debe hacerse según el orden de mayor a menor o viceversa.

Si se considera la variable ordinal grado de escolaridad, por ejemplo, los compartimientos deben disponerse según el orden: sin educación, educación primaria, educación secundaria, educación universitaria.

Salvando las diferencias propias de la medición, la propiedad-espacio de las variables intervalos y de razones pueden considerarse conjuntamente. La cantidad de posiciones disponibles para las unidades la da el refinamiento del instrumento con que se procede a efectuar la medición. Si se mide con un cronómetro que sólo marca hasta décimas de

segundo, el tiempo que demora un individuo en hacer una tarea sólo podrá ser medido con una precisión de décimas de segundo, y si la tarea tiene una duración máxima de un minuto, las posiciones disponibles para los sujetos puestos a prueba serán 600 (60 segundos $\times$ 10 décimas). Se obtiene así una propiedad-espacio de 600 posiciones. Si el cronómetro marca hasta quintos de segundo, la propiedad-espacio resultante para efectuar la medición dispondrá de 1200 posiciones.

Es conveniente recordar que las variables de los niveles más altos de medición pueden tratarse como variables de nivel inferior, pero no a la inversa. Una variable de razones (ingreso) puede tratarse como variable ordinal, pero una variable ordinal (grado de escolaridad) no puede tratarse como de razones.

Supongamos que se desea poner a prueba la hipótesis: “La disposición de los inmigrantes a adquirir la ciudadanía del país de llegada tiende a variar según el país de origen”. En la recolección de los datos se ha preguntado a los individuos acerca de su nacionalidad de origen. Se obtiene así una amplitud de respuestas que puede referirse a todos los países del mundo. Las posiciones disponibles para cada unidad superan el número de 100. Por lo tanto, resulta un manejo engorroso de los datos y la posibilidad de que en algunas de las posiciones se encuentren menos frecuencias que las necesarias para efectuar la manipulación estadístico-matemática, así como el hecho de que carezca de relevancia analizar la actitud de los inmigrantes de países que tienen frecuencias muy bajas.

Cuando éste es el caso, puede procederse a una *reducción de la propiedad-espacio*, que es una operación consistente en combinar el número de clases, con el fin de obtener un número menor de categorías.

La reducción es una operación general que puede realizarse de distintas maneras. Una de dichas maneras es la *simplificación*. Simplificar una propiedad-espacio es juntar posiciones de una variable ordinal, interval o de razones en un número menor de clases ordenadas jerárquicamente, o las posiciones de una variable nominal en un número menor de clases, o en una dicotomía.

En el ejemplo anterior, se procedería a una simplificación de la propiedad-espacio si se juntaran todos los países de algún continente, manteniendo aislados a los países que merecen una consideración especial. O, a la inversa, se mantuvieran aislados ciertos países y se pusiera al resto en una categoría encabezada como “otros países”.

Para que una simplificación sea correcta, debe atenderse a los siguientes criterios: 1) Debe tener relevancia teórico-empírica. Es decir, la simplificación no debe conducir a perder información que sea importante para la comprobación de la hipótesis. 2) Debe atenderse a las reglas de la clasificación. Es decir, la propiedad-espacio resultante debe contener posiciones mutuamente excluyentes para cada una de las unidades y exhaustivas de toda la variable. 3) Debe proveer de frecuencias suficientes en cada posición para las manipulaciones estadístico-matemáticas.

Como se ha hecho notar anteriormente, lo que es permitido al nivel nominal puede hacerse a cualquier otro nivel, siempre que se conserven los mismos criterios.

Mediante la simplificación, toda propiedad-espacio puede reducirse a dos categorías

ordinales o a una dicotomía nominal. Si el asunto se considera atentamente, se ve que muchas operaciones de codificación (especialmente con las preguntas abiertas sobre salario o sobre edad) no son sino simplificaciones previas al manejo de los datos en cuadros estadísticos.

Supongamos que mediante una pregunta abierta se ha determinado el ingreso de los individuos. Se ha tomado la variable ingreso (una variable de razones) y se ha codificado en intervalos de amplitud 100, obteniéndose la siguiente distribución:

CUADRO 1

<i>Ingreso</i>	<i>Frecuencia</i>	<i>Porcentaje</i>	<i>Porcentaje acumulativo</i>
0-100	7	1.9	—
101-200	30	7.7	9.6
201-300	71	18.1	27.7
301-400	93	23.8	51.5
401-500	58	14.9	66.4
501-600	34	8.7	75.1
601-700	29	7.4	82.5
701-800	8	2.1	84.6
801-900	11	2.8	87.4
901-1000	1	0.3	87.7
1001-1100	3	0.8	88.5
1101-1200	—	—	88.5
1201-1300	1	0.3	88.8
1301-1400	1	0.3	89.1
1401-1500	—	—	89.1
1501 y más	3	0.8	89.9
No se aplica	31	8.0	97.9
NS/NR *	8	2.0	99.9

* NS = No sabe: si el entrevistado no conoce la respuesta.

NR = No responde: si el entrevistado no desea contestar la pregunta.

Esta distribución, ya simplificada en la codificación, todavía proporciona un número muy grande de posiciones a las unidades. Es necesario proceder a su simplificación. Para ello deben tenerse presentes los criterios antes entregados, pero no sólo éstos. Es decir, si la variable en cuestión va a ser analizada por separado, conviene reducirla teniendo en cuenta exclusivamente los criterios antes entregados. Pero si la variable va a ser introducida junto con otras en un índice, es conveniente simplificarla teniendo en cuenta, además de los criterios ya citados, la obtención de una proporción igual de casos en cada una de las posiciones resultantes de la simplificación. No es conveniente simplificar una variable que será introducida en un índice, por ejemplo, dejando un 10% de las unidades en una posición extrema, un 50% en la posición media y el 40% restante en la otra

posición extrema.

En el ejemplo propuesto, si el salario vital mínimo fuera de valor 300, se podría intentar la simplificación en base a las siguientes posiciones:

Hasta un salario vital mínimo	27.7%
Uno y dos salarios vitales	47.4%
Más de 2 salarios vitales	14.8%

Según el número de unidades con que se cuenta y si se analiza aisladamente la variable ingreso, esta simplificación puede ser válida. Pero si la variable va a ser introducida en un índice, conviene dejar las posiciones establecidas de la siguiente forma:

Ingreso bajo (0-300)	27.7 %
Ingreso medio (301-500)	38.7 %
Ingreso alto (501 y más)	23.5 %

que provee la distribución más igualitaria de las frecuencias en el ejemplo dado.

Con esta simplificación se tiene la variable de razones ingreso convertida en una variable ordinal tricotomizada. Para mayores detalles puede consultarse el capítulo II, en la sección correspondiente a índices.

La propiedad-espacio en un sistema de dos variables

Un sistema de dos variables define un espacio de dos dimensiones y, por tanto, puede representarse por un plano. Si las variables son intervalos o de razones, el plano resultante puede considerarse como un plano continuo, con todas sus propiedades estadístico-matemáticas. Si se conocen los valores de una unidad en cada una de las variables, puede determinarse fácilmente cuál es el punto que le corresponde en el plano.

Pero si las variables son nominales u ordinales, el plano resultante ya no es continuo, sino un conjunto ordenado de celdas que representa una combinación de los valores de las dos variables.

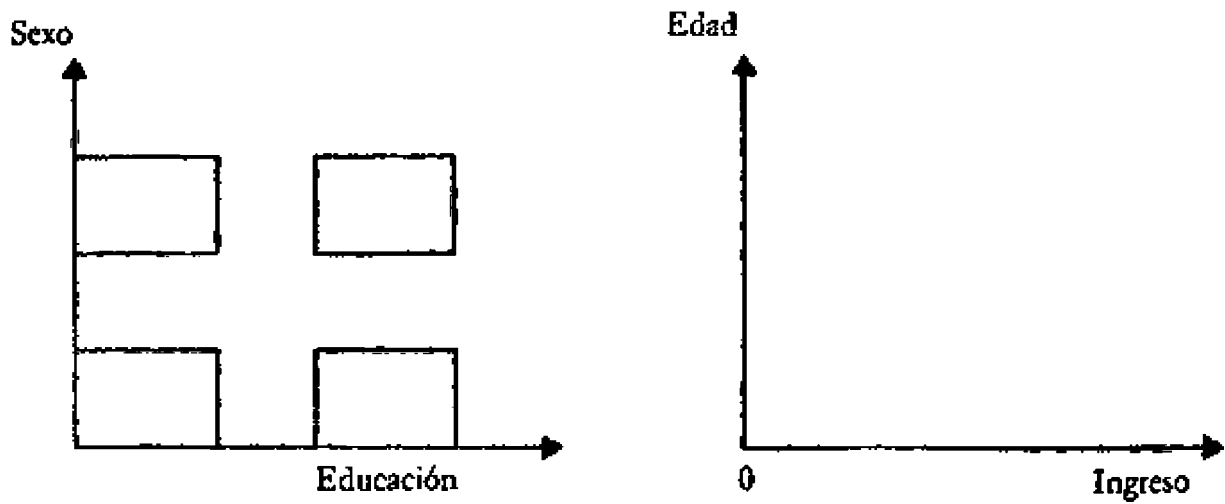
Se puede tomar como ejemplo la combinación de las variables sexo-educación y edad-ingreso. Los planos definidos serían los de la gráfica 1.

En cualquier caso (salvo cuando se combinan variables nominales entre sí o combinadas con variables de otro nivel de medición), la intersección de las dos variables debe producirse en el mismo nivel de ambas variables. No es posible intersectar las variables edad-ingreso en el valor máximo de ingreso y en el mínimo de edad.

El número de posiciones o celdas resultante de combinar dos o más variables reconoce dos situaciones distintas:

1) Cuando las variables tienen, por separado, el mismo número de posiciones, se aplica la expresión p^v , donde p es el número de posiciones y v es el número de variables. Así, si se tienen dos variables dicotomizadas, el número de posiciones definidas es $p^v =$

$2^2 = 4$; si se tienen dos variables tricotomizadas, el número de posiciones definidas es $p^v = 3^2 = 9$.



GRÁFICA 1

2) Cuando las variables tienen, por separado, distinto número de posiciones, se aplica la expresión $p_1 \times p_2 \times \dots \times p_n$ donde p_1 es el número de posiciones de la variable 1, p_2 el número de posiciones de la variable 2, etc. Así, de la combinación de tres variables, una dicotomizada y las otras dos tricotomizadas, las posiciones resultantes son: $p_1 \times p_2 \times p_3 = 2 \times 3 \times 3 = 18$ posiciones.

El número de posiciones que se definen en una propiedad-espacio es importante para estimar el número de unidades necesario para realizar una investigación. En principio deben presupuestarse alrededor de 30 unidades por posición. Así, para una propiedad-espacio definida por dos variables tricotomizadas, debieran preverse 270 unidades, resultantes de la aplicación de la expresión:

$$(p^v \times 30) = 3^2 \times 30 = 9 \times 30 = 270$$

Es costumbre denominar las posiciones de una propiedad-espacio a partir de la posición superior izquierda, procediendo horizontalmente, a saber:

a	b	c
d	e	f
g	h	i

Supongamos que con el fin de probar la hipótesis I: “Mientras más bajo el nivel económico social (NES),² más favorable tenderá a ser la actitud hacia el partido político X y viceversa”, se dispone de las siguientes distribuciones:

<i>Actitud hacia el partido X</i>	<i>Frecuencia</i>
Muy favorable	2 100
Favorable	1 290
Indiferente	3 580
Desfavorable	1 030
Muy desfavorable	1 000
	9 000

<i>Estrato de NES</i>	<i>Frecuencia</i>
Bajo-bajo	2 200
Bajo-alto	1 950
Mediano-bajo	1 700
Mediano-alto	1 550
Alto-bajo	850
Alto-alto	750
	9 000

Con estas dos variables y sus respectivas posiciones, se tiene definida una propiedad-espacio de $p_1 \times p_2$ posiciones, en este caso: 30 posiciones, que darían un cuadro como el siguiente:

CUADRO 2

<i>NES</i> <i>Actitud</i>	<i>Bajo-Bajo</i>	<i>Bajo-Alto</i>	<i>Medio-Bajo</i>	<i>Medio-Alto</i>	<i>Alto-Bajo</i>	<i>Alto-Alto</i>	<i>Total</i>
Muy favorable	a 1000	b 1000	c 50	d 50	e —	f —	2100
Favorable	g 400	h 390	i 200	j 200	k 50	l 50	1290
Indiferente	m 600	n 410	o 1150	p 1000	q 300	r 120	3580
Desfavorable	s 200	t 150	u 200	v 200	w 100	x 180	1030
Muy desfavorable	y —	z —	a' 100	b' 100	c' 400	d' 400	1000
<i>Total</i>	2200	1950	1700	1550	850	750	9000

Conviene hacer notar algunos aspectos: 1) La intersección de las dos variables no ha sido hecha en los valores más bajos de ambas variables. Esto es conveniente cuando la hipótesis sostiene una relación inversa: mientras más de una variable, menos de la otra.

El propósito de efectuar la intersección tal como está en el cuadro es observar el comportamiento de las frecuencias en la diagonal principal (la que atraviesa el cuadro desde la posición superior izquierda hasta la posición inferior derecha). 2) El número de posiciones es sumamente elevado. 3) La variable que se considera más independiente se coloca en la parte superior del cuadro. De ahora en adelante llamaremos x a la variable independiente y y a la variable dependiente. 4) Que los totales son iguales a las distribuciones de cada variable, como no podía ser de otra manera. Estos totales (tanto los del extremo derecho como los del extremo inferior) se denominan “marginales”. En el ejemplo, se tienen marginales de x y marginales de y .

Para proceder a reducir esta propiedad-espacio, se combinan posiciones de ambas variables. La variable x se agrupa en una tricotomía con valores bajo, medio y alto; la variable y , en una tricotomía con valores favorable, indiferente, desfavorable. La propiedad-espacio, así definida, tendrá sólo nueve posiciones y el cuadro resultante será el siguiente:

CUADRO 3

NES Actitud	<i>Bajo</i>	<i>Medio</i>	<i>Alto</i>	<i>Total</i>
Desfavorable	$a+b+g+h$ a 2790	$c+d+i+j$ b 500	$e+f+k+l$ c 100	$a+b+c$ 3390
Indiferente	$m+n$ d 1010	$n+\tilde{n}$ e 2150	$o+p$ f 420	$d+e+f$ 3580
Favorable	$q+v+w+x$ g 350	$s+t+y+z$ h 600	$u+v+a'+b'$ i 1080	$g+h+i$ 2030
Total	$a+d+g$ 4150	$b+e+h$ 3250	$c+f+i$ 1600	9000

Pero los análisis sociológicos a partir de dos variables sólo proporcionan información acerca de la asociación entre ellas, sin hacer un gran aporte de carácter explicativo. Para esto es menester agregar por lo menos una tercera variable, para ver si una vez introducida ésta, la asociación entre las dos primeras variables se mantiene, se especifica, se refuerza o desaparece.

Antes de considerar la propiedad-espacio resultante de la introducción de una tercera variable, es conveniente analizar una especie muy útil de reducción de la propiedad-espacio, la llamada *reducción pragmática*.³ Siguiendo con el ejemplo propuesto,

supongamos que se pide poner a prueba la hipótesis II: “La movilidad ascendente tiende a hacer menos favorable la actitud hacia el partido político X y viceversa”. Para comprobar esta hipótesis, se proporciona una tercera variable: NES del padre de individuo-unidad, cuya distribución es la siguiente:

<i>NES del padre</i>	<i>Frecuencias</i>
Estrato bajo-bajo	2750
Estrato bajo-alto	2450
Estrato medio-bajo	1650
Estrato medio-alto	1350
Estrato alto-bajo	500
Estrato alto-alto	300
Total	9 000

Si se procede a la reducción de esta variable por el procedimiento de la simplificación, se obtiene la siguiente distribución:

<i>NES del padre</i>	<i>Frecuencias</i>
Estrato bajo	5200
Estrato medio	3000
Estrato alto	800
Total	9 000

CUADRO 4

NES del padre NES del hijo	Bajo	Medio	Alto	Total
Bajo	a 3700	b 400	c 50	4150
Medio	d 1000	e 2100	f 150	3250
Alto	g 500	h 500	i 600	1600
Total	5200	3000	800	9000

Al formar la propiedad-espacio correspondiente a las variables NES del hijo y NES del padre (2 variables tricotomizadas), quedan definidas 9 posiciones con las frecuencias que se observan en el cuadro 4.

Si se observa atentamente esta propiedad-espacio con 9 posiciones en relación con la movilidad social, se ve claramente que en las posiciones ubicadas sobre la diagonal principal (celdas *a*, *e*, *i*) se encuentran los individuos inmóviles, es decir, aquellos que están en el mismo estrato que su padre. Si se observan las celdas *d*, *g*, *h*, se ve que en ellas se encuentran los individuos con movilidad ascendente, es decir, aquellos individuos de NES medio con padres de NES bajo y los individuos de NES alto con padre de NES medio o bajo. Asimismo, al observar las celdas *b*, *c* y *f*, se obtiene a los individuos móviles descendentes, es decir, aquellos individuos que, teniendo padre de NES alto o medio, están en NES bajo, y los individuos que, con padre de NES alto, están en NES medio.

De tal modo, es posible efectuar la reducción pragmática de la propiedad-espacio original de 9 posiciones a una propiedad-espacio de sólo tres posiciones, a saber:

<i>Movilidad</i>	<i>Frecuencias</i>	<i>Modo de obtenerla</i>
Descendentes	600	Suma de celdas $b + c + f$
Inmóviles	6400	Suma de celdas $a + e + i$
Ascendentes	2000	Suma de celdas $d + g + h$
Total	9000	

Debe notarse que en la reducción pragmática:

1) Normalmente desaparecen las variables que definían la propiedad-espacio original y se obtiene una nueva variable. En el ejemplo, ya no se tiene NES del padre o NES del hijo, sino una nueva variable: dirección de la movilidad social. Se ha producido así una reducción de la propiedad-espacio, no en virtud de una simplificación del número de posiciones de las variables, sino en virtud de una transformación de las dimensiones del espacio dado que, de un espacio de dos dimensiones, ha pasado a ser un espacio de una dimensión. (Diríase que un plano se ha transformado en una línea.)

2) La reducción pragmática se justifica en atención al significado teórico de las posiciones definidas por la propiedad-espacio original y no por los valores numéricos de estas posiciones. (Esto se verá más claramente al considerar la reducción numérica de la propiedad-espacio.)

3) La reducción pragmática es una operación sumamente útil cuando se efectúan estudios secuenciales, es decir, cuando se analizan los distintos valores que va teniendo una unidad en una misma variable a lo largo del tiempo. Esto permite proceder a la reducción pragmática en base a categorías tales como constancia-inconstancia, estabilidad-inestabilidad, etc., ya sea el diseño del estudio en un solo tiempo o un diseño tipo panel.

Así, puede decirse que la *reducción pragmática* es una combinación de las posiciones de una propiedad-espacio en busca de significaciones teóricas nuevas, a partir del significado teórico de las posiciones y no de sus valores numéricos.

Retomando el ejemplo, la propiedad-espacio definida por las variables dirección de la movilidad y actitud hacia el partido X, necesaria para decidir si la hipótesis puede ser aceptada o debe ser rechazada, sería la siguiente:

CUADRO 5

<i>Dirección de la movilidad</i> <i>Actitud</i>	<i>Descendente</i>	<i>Inmóvil</i>	<i>Ascendente</i>	<i>Total</i>
Favorable	290	2800	300	3390
Indiferente	210	2570	800	3580
Desfavorable	100	1030	900	2030
Total	600	6400	2000	9000

Debe notarse que los grupos experimentales están definidos por la dirección de su

movilidad. La lógica del razonamiento es la siguiente: formados tres grupos distintos según la dirección de la movilidad, se somete a todos ellos a un mismo estímulo. Si las respuestas al estímulo son distintas en los diferentes grupos y se producen en la dirección prevista, la hipótesis puede aceptarse. Si las respuestas al mismo estímulo no muestran diferencias en los distintos grupos o la muestran en una dirección no prevista en la hipótesis, ésta debe ser rechazada.

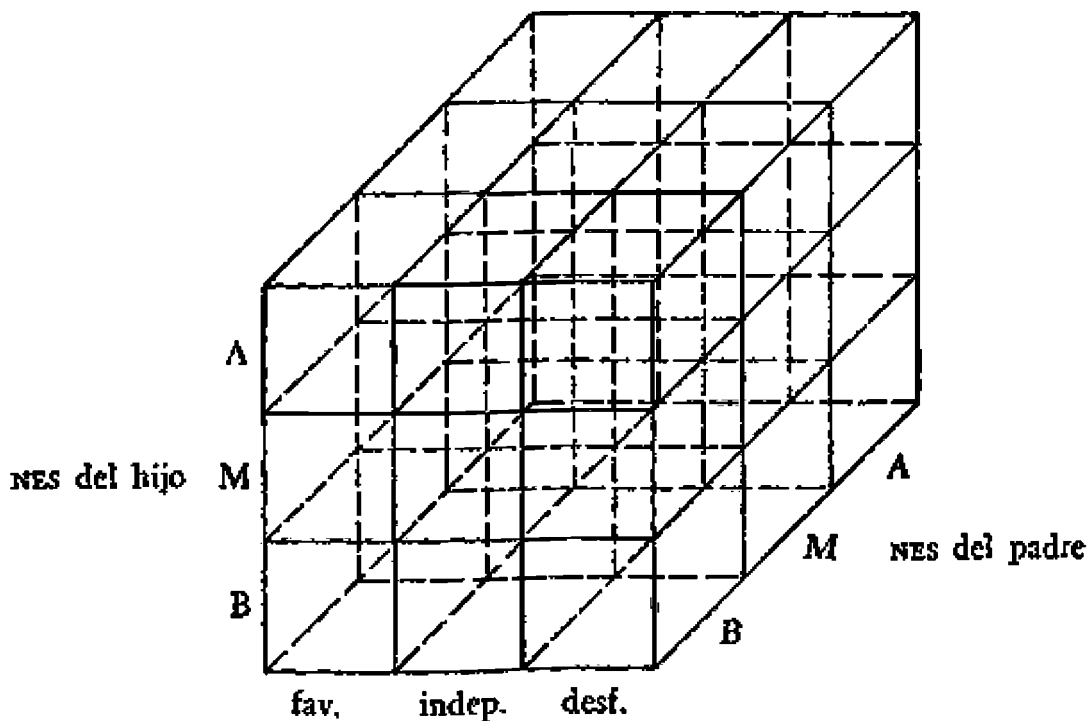
La propiedad-espacio en un sistema de tres variables

Un sistema de tres variables define un espacio de tres dimensiones y, por tanto, puede representarse por un cubo.

Si las tres variables tienen el mismo número de posiciones, el sistema tendrá tantas posiciones como resulten de la aplicación de la expresión p^v , antes considerada.

Si las variables tienen distinto número de posiciones, el sistema tendrá tantas como resulten de multiplicar la cantidad de posiciones de todas las variables, aplicando la expresión $p_1 \times p_2 \times \dots \times p_n$, que también ya ha sido considerada.

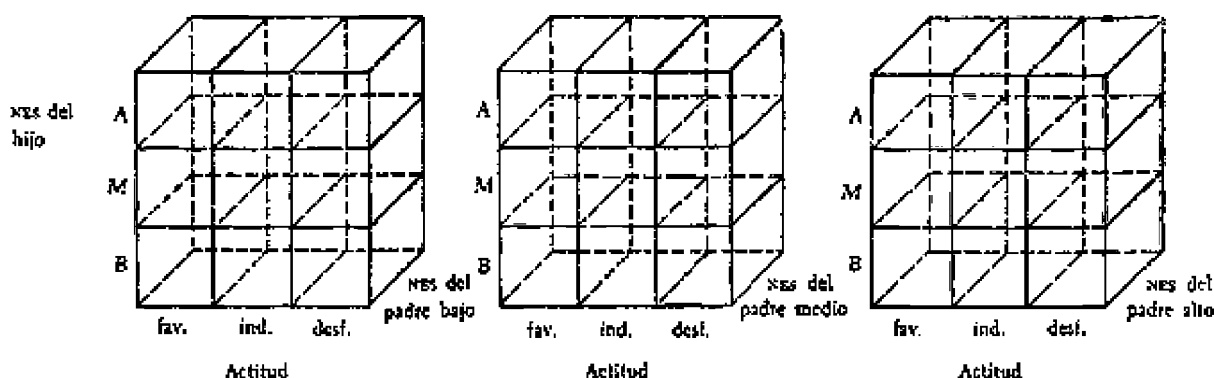
Mediante un ejemplo y retomando la representación del cubo, se verá más claramente la consecuencia (por ahora gráfica) de introducir una tercera variable. Si a las variables “Actitud hacia el partido X” y “NES del individuo”, se agrega el “NES del padre”, cada una de ellas tricotomizadas, el cubo resultante sería el siguiente:



Actitud

GRÁFICA 2

Como la representación gráfica global es decididamente complicada, se pueden presentar tres cortes en el cubo, cada uno de los cuales corresponde a la tercera variable introducida. La representación gráfica, en consecuencia, sería la siguiente:



GRÁFICA 3

Nótese que en cada sección del cubo se mantienen todas las posiciones de las variables NES del hijo y actitud hacia el partido político X, y que cada sección corresponde a una posición de la variable “NES del padre”. Es el simple recurso de mantener el espacio bidimensional colocando en el título la tercera variable introducida, lo que permite una representación más manejable, como se muestra en el cuadro 6.

CUADRO 6

Actitud hacia el partido político "X"	Padre de NES bajo NES del hijo				Padre de NES medio NES del hijo				Padre de NES alto NES del hijo			
	Bajo	Medio	Alto	Total	Bajo	Medio	Alto	Total	Bajo	Medio	Alto	Total
Favorable	2600	200	50	2850	150	200	50	400	40	100	—	140
Indiferente	850	500	150	1500	150	1600	150	1900	10	50	120	180
Desfavorable	250	500	300	850	300	300	500	700	—	—	480	480
Total	3700	1000	500	5200	400	2100	500	3000	50	150	680	880

Éste es el problema gráfico de introducir una tercera variable. Pero la introducción de una tercera variable es inseparable de las finalidades del análisis. Como se señaló antes, los análisis a partir de dos variables sólo expresan la asociación entre éstas, sin hacer un gran aporte de carácter explicativo. Sólo la introducción de una tercera variable permite decidir, dentro del sistema de variables que se considera, si la relación entre las dos primeras variables se mantiene, se refuerza, se especifica o desaparece.

Supongamos que se quiere poner a prueba la hipótesis III: “Cualquiera que sea el nivel económico-social del padre, se mantiene la H_I ”.

Retomando nuevamente los problemas de lógica subyacente, cabe decir que, con la introducción de una tercera variable, se multiplica el número de grupos experimentales que se tenían en el análisis de dos variables por la cantidad de posiciones que tiene la tercera variable introducida. En el ejemplo, de los tres grupos experimentales iniciales

según el NES del individuo se pasa a tener nueve grupos experimentales (padre de NES bajo, hijo de NES bajo; padre de NES bajo, hijo de NES medio; padre de NES bajo, hijo de NES alto; padre de NES medio, hijo de NES bajo; padre de NES medio, hijo de NES medio; etc.). De esta manera, se está en condiciones de decidir sobre la hipótesis III, comparando los grupos experimentales, en base a las siguientes preguntas:

1) Para cada nivel económico-social del padre, ¿se mantiene la relación primitiva entre NES del hijo y actitud hacia el partido político X?

2) Para cada nivel económico-social del hijo, ¿se mantiene su relación con la actitud hacia el partido político X, si se consideran los distintos niveles económico-sociales de los padres?

Las respuestas a estas preguntas, no sólo las respuestas afirmativas o negativas, sino las tendientes a explicar por qué, cómo y cuándo las relaciones se mantienen, se refuerzan, se especifican o desaparecen, son básicas en el análisis de los datos.

Es conveniente considerar ahora un problema menos profundo que el del análisis, pero de mucha utilidad en el manejo de los datos. Como se ha hecho hasta ahora, supóngase que se pide la demostración de la hipótesis IV: “Cualquiera que sea el grado de desarrollo alcanzado por el país latinoamericano, se mantiene la hipótesis III”, y para determinar el grado de desarrollo de los países latinoamericanos, se proporcionan los datos del cuadro 7.

Si se pretendiera manejar la propiedad-espacio definida por las posiciones en cada variable, el número de posiciones del sistema de tres variables sería elevado ($20 \times 20 \times 20$). Por otra parte, interesa obtener un valor único. Ya se ha visto que esto puede lograrse (cuando se han dado ciertas condiciones) por medio de un índice sumatorio. Para la construcción del índice, se procede a ordenar de mayor a menor (o viceversa) los valores registrados en cada una de las variables. Las columnas 5, 6 y 7 del cuadro 7 registran este ordenamiento, prescindiendo del país al que corresponde la hilera.

Analizando cada una de las columnas 5, 6 y 7 por separado, se procede a simplificar las veinte posiciones dadas a sólo tres (siguiendo las líneas punteadas), según las consideraciones ya efectuadas al estudiar los índices sumatorios y la propiedad-espacio de una sola variable. La propiedad-espacio resultante (tres variables tricotomizadas) tiene $p^v = 27$ posiciones, a cada una de las cuales corresponde un valor en el índice sumatorio el que, a su vez, tiene un puntaje mínimo 0 y un puntaje máximo 6.

CUADRO 7

(1)	(2)	(3)	(4)	(5)	(6)	(7)
<i>Países</i>	<i>% de po- blac. acti- va en ac- tividades fabriles propiam. dichas</i>	<i>% de po- blac. en ciud. de 20000 habit. y más</i>	<i>% de po- blac. que sabe leer</i>	<i>Orden col. 2</i>	<i>Orden col. 3</i>	<i>Orden col. 4</i>
Argentina	13.5	28	87	13.5	28	95
Bolivia	2.8	7	32	13.0	28	87
Brasil	6.2	13	49	9.1	21	80
Colombia	4.3	12	63	8.0	21	80
Costa Rica	4.8	14	80	6.7	16	78
Cuba	8.0	21	78	6.5	15	70
Chile	9.1	21	80	6.2	14	66
Ecuador	3.5	10	56	4.8	14	63
El Salvador	4.7	9	57	4.7	13	59
Guatemala	3.2	6	20	4.7	12	57
Haití	1.4	2	11	4.3	12	56
Honduras	1.6	4	35	3.5	12	52
México	6.7	14	59	3.3	10	49
Nicaragua	2.0	4	38	3.2	10	43
Panamá	3.0	15	70	3.0	9	42
Paraguay	3.3	12	66	3.0	7	38
Perú	4.7	10	42	2.8	6	35
Rep. Dominicana	3.0	12	43	2.0	4	32
Uruguay	13.0	28	95	1.6	4	20
Venezuela	6.5	16	52	1.4	2	11

Como ya se ha visto con los índices sumatorios, si se consideran los puntajes extremos (0 y 6), no se presenta ningún problema especial. Esos puntajes pueden alcanzarse de una sola manera posible: teniendo el mismo puntaje en cada una de las variables. Con los restantes valores (1 a 5) se presentan problemas, pues éstos ya pueden obtenerse con distintas combinaciones de valores de las tres variables. El problema tiene una solución práctica:

1) Se consideran "tipos puros" todas aquellas unidades que obtienen el mismo puntaje en cada una de las tres variables (bajo-bajo-bajo, alto-alto-alto, medio-medio-medio).

CUADRO 8

Ocupación: <i>fabril</i>		Baja (1,4 - 3,3) Valor: 0			Media (3,5 - 6,2) Valor: 1			Alta (6,5 - 13,5) Valor: 2		
Alfabetización	Urbanización	Baja (2 - 10) Valor 0	Media (12 - 14) Valor 1	Alta (15 - 28) Valor 2	Baja (2 - 10) Valor 0	Media (12 - 14) Valor 1	Alta (15 - 28) Valor 2	Baja (2 - 10) Valor 0	Media (12 - 14) Valor 1	Alta (15 - 28) Valor 2
Baja (11.43) Valor 0	0 Bolivia, Guatemala, Haití, Honduras, Nicaragua	1 R. Dominicana	2	1 Perú	2	3	2	3	4	
Media (49.66) Valor 1	1	2	3	2 Ecuador El Salvador	3	4	3	4	5	Venezuela
Alta (70.95) Valor 2	2	3	4 Panamá	3	4 Costa Rica	5	4	5	6	Argentina, Cuba, Chile, Uruguay

CUADRO 9

NEA del padre	Actitud hacia el partido político X	Nivel de desarrollo Bajo				Nivel de desarrollo Medio				Nivel de desarrollo Alto				
		NEA DEL HIJO				NEA DEL HIJO				NEA DEL HIJO				
		Bajo	Medio	Alto	Total	Bajo	Medio	Alto	Total	Bajo	Medio	Alto	Total	TOTAL
Bajo	Favorable	1000	100	50	1150	900	50	—	950	700	50	—	750	2850
	Indiferente	300	200	50	550	300	200	50	550	250	100	50	400	1500
	Desfavorable	100	100	100	300	100	50	50	200	50	150	150	350	850
	Total	1400	400	200	2000	1300	300	100	1700	1000	300	200	1500	3200
Medio	Favorable	50	100	—	150	100	50	50	200	—	50	—	50	400
	Indiferente	50	300	50	400	50	500	50	600	50	800	50	900	1900
	Desfavorable	—	100	50	150	50	50	100	200	50	150	150	350	700
	Total	100	500	100	700	200	600	200	1000	100	1000	200	1300	3000
Alto	Favorable	—	50	—	50	40	50	—	90	—	40	—	40	140
	Indiferente	—	20	50	70	10	20	50	80	—	10	20	30	180
	Desfavorable	—	—	200	200	—	—	150	150	—	—	150	150	480
	Total	—	50	250	300	50	50	200	300	—	50	150	200	500
Total		1500	950	550	3000	1550	950	500	3000	1100	1350	550	3000	9000

2) Los puntajes restantes se adjudican de la siguiente manera: a) Dos puntajes iguales en dos variables y uno inmediato en la tercera, se otorga el puntaje obtenido en las dos variables (bajo-bajo-medio, se adjudica a valor bajo; alto-alto-medio, a valor alto; medio-medio-bajo o medio-medio-alto, ambos a valor medio). b) Dos puntajes extremos en dos variables y el puntaje opuesto en la tercera: se adjudica a valor medio (bajo-bajo-alto: a valor medio; alto-alto-bajo: a valor medio).

Con este procedimiento se obtienen tres valores: 1) el valor alto obtenido según a y $b1$; 2) el valor medio: obtenido según a , $b1$ y $b2$, y 3) el valor bajo, obtenido según a y $b1$.

Esta reducción, que se basa en los valores numéricos acordados a las distintas posiciones de las variables, es la llamada *reducción numérica* (compárese con la reducción pragmática para retener mejor la diferencia entre ambas reducciones).

Al obtener esta variable y volviendo al ejemplo, se tiene una propiedad-espacio definida por cuatro variables: grado de desarrollo del país, NES del padre, NES del hijo y actitud hacia el partido político X. Es decir, que la propiedad-espacio definida por estas cuatro variables tricotomizadas (según la expresión p^v) proporciona 81 posiciones, tal como se observa en el cuadro.

Un cuadro de estas dimensiones es, lisa y llanamente, ilegible. Las soluciones de presentación, que, por otra parte, son soluciones que facilitan el análisis de los datos, pueden ser dos: 1) La solución ya considerada de poner en el título una de las variables. 2) La solución de presentar parcialmente los datos, efectuando el análisis de cada parte por separado.

La substrucción

Es frecuente encontrar en la literatura sociológica concepciones como la que presentan Gerth y Mills sobre “la orientación hacia los controles sociales”: ⁴

i) Existe el tipo de hombre que estima o afirma una determinada norma, y en su conducta se sujeta a su convicción. *ii)* Existe también el conformista y el oportunista falso. Externamente se ajusta al código, internamente no lo acepta. *iii)* Un tercer tipo está representado por la persona que acepta verbalmente un código, pero que, en la conducta, se desvía de él. *iv)* Este tipo representa al desviado consistente, al disconforme en palabras y en hechos.

Como hacen notar los autores, estos cuatro tipos se mueven en un sistema de dos variables: “la actitud hacia el ideal o la norma” y “la conducta con relación a la norma o ideal”, cada una de ellas con dos valores: positivo y negativo. Así, detrás de la “tipología” presentada se encuentra una propiedad-espacio definida por un sistema de dos variables dicotomizadas que, tal como hacen los autores, puede presentarse así:

<i>Actitud hacia el ideal o la norma</i>			
		+	—
Conducta con relación a la norma o ideal	+	Tipo I	Tipo II
	—	Tipo III	Tipo IV

Pero no siempre los autores tienen el cuidado de presentar (o no se han percatado de) las variables subyacentes a sus conceptualizaciones y el panorama no se presenta claro.

La substrucción, así, es el procedimiento por el cual se intenta reconstituir el sistema de variables que definen la propiedad-espacio subyacente a uno o varios conceptos.

Si se analiza la conocida definición de anomia como “ausencia de normas que regulan los medios y caminos que llevan a una meta socialmente aceptada”,⁵ puede verse que, subyaciendo a la definición, hay dos variables: metas socialmente aceptadas y normas sobre medios y caminos para lograr las metas. Cada una de esas variables puede pensarse con dos posiciones: existencia-inexistencia. Con ello se puede reconstituir la propiedad-espacio definida:

		Metas socialmente aceptadas	
		Existen	No existen
Normas sobre caminos o medios	Existen		
Logro de las metas	No existen	Anomia	

De tal modo, la substrucción es un procedimiento inverso a la reducción pragmática o a la clasificación articulada, que permite reconstituir: 1) Las variables incluidas en los conceptos o en la clasificación articulada. 2) La propiedad-espacio definida por el sistema de variables. 3) La reducción pragmática aplicada.

Algunos autores (Barton entre ellos)⁶ consideran que la substrucción es aplicable especialmente al tratar términos que designan individuos, o situaciones o relaciones “típicas”, o al considerar “tipologías”. Esto es acertado si no se confunden los términos “típico” o “tipológico” con los “tipos ideales” weberianos. El procedimiento es útil y válido cuando se piensa que el autor ha llegado a la formulación de los tipos o las tipologías por medio de una clasificación articulada o las combinaciones lógicas de los distintos valores de diferentes variables. Pero no se justifica en el caso de los tipos ideales que resulten o no de una clasificación articulada o de una combinación lógica de valores de variables; se basan metodológicamente en la comprensión y no en alguna forma de reduccionismo lógico-matemático.

B) RAZONES, TASAS, PROPORCIONES Y PORCENTAJES

Considerados ya los aspectos más importantes de la propiedad-espacio y con ello el primer tema enunciado al comenzar su tratamiento, queda por analizar el segundo tema, que hace referencia al tratamiento de los datos, utilizando las herramientas que proveen las matemáticas y la estadística.

En este momento sólo cabe hacer el planteo introductorio al manejo más elemental de los datos por medio de razones, tasas, proporciones y porcentajes.

Previamente, conviene tener a la vista cuáles son las operaciones realizadas hasta

llegar a este momento de la investigación. A partir de una hipótesis, se han recogido y codificado los datos y se ha definido la propiedad-espacio necesaria para disponer los datos de modo que permitan decidir sobre la aceptación o rechazo de la hipótesis.

A continuación, lo obvio es el conteo. Es decir, la determinación de cuántas unidades entran en cada posición de la propiedad-espacio. Esto proporciona el número absoluto de unidades ubicadas en cada posición. Pero el número absoluto de unidades en cada posición es la cantidad de unidades de un grupo experimental que ha dado una respuesta determinada a un estímulo. Estos números absolutos son mayores, menores o iguales que los números absolutos que hay en otras posiciones de la propiedad-espacio. El hecho de que el número absoluto de una posición sea mayor, menor o igual que los de otras posiciones no permite extraer ninguna conclusión sobre el comportamiento de las unidades, ya que es muy distinto el significado de 100 unidades en un grupo de 200 que el de 1000 unidades en un grupo de 5 mil.

Para poder extraer conclusiones es necesario poder comparar y para poder comparar es necesario recurrir a números relativos: números puestos en relación con otros números del mismo grupo experimental. Los números que están en cada posición de un grupo experimental pueden ponerse en relación: *a)* Con los valores absolutos de cada una de las otras posiciones del mismo grupo experimental. Estos números relativos son las razones y tasas. *b)* Con el total de las unidades del grupo experimental. Con este recurso, todos los grupos experimentales se igualan en el valor 1. Los números relativos así obtenidos son las proporciones. *c)* Con el total de unidades del grupo experimental, pero igualado con los otros grupos experimentales en el valor 100. Estos números relativos son los llamados porcentajes.

Todo lo expuesto hasta el momento se verá más claramente con un ejemplo. Supóngase la hipótesis: “La movilidad ascendente tiende a hacer menos favorable la actitud hacia el partido político, y viceversa”.

El cuadro de que se dispone es el siguiente:

Antes de considerar los números relativos, debe notarse: *1)* Que se tienen tres grupos experimentales: individuos con movilidad descendente, individuos inmóviles, individuos con movilidad ascendente, cuyos totales respectivos son 600, 6400 y 2000. *2)* Que cada uno de esos grupos ha respondido en forma diferencial al mismo estímulo: la pregunta o el conjunto de preguntas para descubrir su actitud hacia el partido político X. *3)* Que los números absolutos no permiten extraer conclusiones. Las 300 unidades de la posición o celda *c* son más que los 290 de la posición o celda *a*. En efecto, ¿puede decirse que los individuos con movilidad ascendente tienen una actitud más favorable hacia el partido político X que los individuos con movilidad descendente? Obviamente no, por la simple razón de que 300 en 2000 es mucho menos (en números relativos) que 290 en 600.

CUADRO 10

<i>Actitud hacia el partido político X</i>	<i>Movilidad descendente</i>		<i>Movilidad nula</i>		<i>Movilidad ascendente</i>		<i>Total</i>
Favorable	290	48%	2800	44%	300	15%	3390
Indiferente	210	35%	2570	40%	800	40%	3580
Desfavorable	100	17%	1030	16%	900	45%	2030
Total	600	100%	6400	100%	2000	100%	9000

Razones

La razón de un número A con respecto a otro número B se define como A/B o $A:B$. Si B fuera el total del grupo, no se trataría de una razón, sino de una proporción.

Con respecto a las razones, debe notarse que:

1) Una razón puede ser menor, igual o mayor que 1. En el ejemplo, la razón de los individuos con movilidad descendente con actitud favorable hacia el partido político X a los individuos con movilidad descendente con actitud desfavorable hacia el partido político X es $290:100 = 2.9$, es decir: $2.9:1$. La inversa, que también es válida, es $100:290 = 1:2.9$, es decir, 0.34 .

2) Los números A y B pueden expresar cantidades distintas cualitativamente. Se puede pensar en una razón de hombres a mujeres, de peras a manzanas, etcétera.

3) Debe tratarse de simplificar la expresión, eliminando del numerador y del denominador los factores comunes.

4) Es conveniente expresar la razón reduciendo el denominador a la unidad. Si se toman los individuos inmóviles y se obtiene la razón de favorables a desfavorables, se tiene $2.800:1.030$. Si esta razón se expresa como $2.80:1.03$, es difícil su interpretación. Pero si se expresa como $2.71:1$, la interpretación se facilita.

5) Si las posiciones son dos, la razón puede calcularse a partir de la proporción, y viceversa. Si se toma la razón de 290 a 100 (prescindiendo de los 210 indiferentes con movilidad descendente), el total sería 390. La razón de $290:390$ y de $100:390$ se expresaría aproximadamente como $3:4$ y $1:4$. Pero, $290+100 = 390$ y la proporción sería $\frac{290}{390} + \frac{100}{390} = 0.74+0.26$, lo que es lo mismo que decir $3:4$ y $1:4$, aproximadamente.

6) El cálculo de las razones tiene un grave inconveniente de legibilidad: a medida que aumenta el número de categorías dentro de cada grupo experimental el manejo de las razones se vuelve engorroso. Si dentro de un grupo experimental hay tres posiciones posibles, el número de razones que se puede obtener resulta de la combinatoria de posiciones, según la expresión $p!$ (factorial de p .) Si en el ejemplo se toma a los individuos con movilidad ascendente, las razones posibles con tres categorías serían: $300:900$; $300:800$; $800:900$; $900:800$; $900:300$; $800:900$. Es decir, $3! = 6$. Si se tuvieran cinco posiciones en un grupo experimental, las razones posibles serían $5! = 120$. Esto es lo que determina la escasa utilización de las razones en análisis sociológicos.

En el ejemplo propuesto, tomando la razón de individuos con actitud favorable a individuos con actitud desfavorable, se tendría:

Movilidad ascendente

2.9

Inmóvil

2.71

Movilidad descendente

0.15

con lo que se tendría demostrada la primera parte de la hipótesis.

Tasas

Las tasas no son sino una clase de razón que, por tener un denominador muy grande que haría impracticable, difícil o ilegible el resultado, se obtiene multiplicando el numerador por una constante $k = 100, 1000, 10000$, etcétera.

Hay tasas ya establecidas en que se conoce el valor de la constante (k). De esta naturaleza son la mayoría de las tasas utilizadas en demografía. Pero si la tasa que se utiliza no es del tipo antes enunciado, hay que dejar constancia expresa del valor que se acuerde a la constante (k). Por ejemplo, si se calcula la tasa de criminalidad de localidades de menos de 100 mil habitantes, la constante k deberá tener un valor de 10000 o de 100000; pero si se calcula la misma tasa de criminalidad para ciudades de más de 1000 000 de habitantes, la constante k deberá valer 100000 o 1000000.

Una tasa muy importante es la *tasa de crecimiento*, que se obtiene tomando los valores de una unidad en una variable en dos tiempos distintos (el tiempo inicial $-T1$ y el tiempo final $-T2$). Esta tasa de crecimiento se calcula aplicando la siguiente expresión: $\frac{T2 - T1}{T1}$. Esto es, se mide cuánto se ha incrementado la unidad desde el tiempo inicial al tiempo final y se divide por el tiempo inicial. Esta tasa es de aplicación muy frecuente en estudios sobre urbanización, desarrollo económico, etc. El problema subsistente es el de la homogeneidad de los periodos, ya que las tasas de crecimiento de dos unidades no serían comparables si se toma un periodo para una unidad (supóngase 5 años) y otro periodo para la otra (supóngase 10 años). La tasa de crecimiento puede ser positiva o negativa.

Proporciones

Una proporción se obtiene dividiendo el número de casos en cada posición de una variable (o de un grupo experimental) entre el total de casos de esa variable (o de un grupo experimental).

En el ejemplo propuesto y para los móviles ascendentes, se tendrían las proporciones siguientes: 300/2000; 800/2000 y 900/2000.

Respecto a las proporciones debe notarse que:

- 1) La proporción es lo que en estadística se llaman frecuencias relativas.
- 2) Las variables o grupos experimentales de los que se calculan las proporciones deben estar correctamente clasificados (en forma exhaustiva y mutuamente excluyente).
- 3) La suma de los numeradores de las proporciones de cada grupo es igual al

denominador o, lo que es lo mismo, al total. En el ejemplo: $300 + 800 + 900 = 2000$.

4) Cada proporción es siempre menor o igual a 1.

5) La suma de todas las proporciones de cada variable o de cada grupo experimental es igual a 1. *Debe notarse* que, con este procedimiento, los grupos experimentales o variables se igualan a 1. Este procedimiento de igualar los totales en una cifra se llama “normalización” o “estandarización”.

6) Es importante tener en cuenta sobre qué total se calculan las proporciones. Si se vuelve al ejemplo, se podría disponer los datos, según la dirección de la movilidad, en la siguiente forma:

Móviles		2600	0.29	—	—
Ascendentes:	2000		—	0.22	0.77
Descendentes:	600		—	0.07	0.23
Inmóviles		6400	0.71	0.71	—
		<u>9000</u>	<u>1.00</u>	<u>1.00</u>	<u>1.00</u>

Según este ejemplo, la proporción de móviles e inmóviles se calcula sobre un total de 9000 unidades. $2600/9000$ y $6400/9000 = 0.29 + 0.71 = 1.00$. Las proporciones de móviles descendentes, ascendentes e inmóviles también se calculan sobre un total de 9 000: $2000/9000$, $600/9000$ y $6400/9000$, es decir: $0.22 + 0.07 + 0.71 = 1.00$. Pero las proporciones de móviles ascendentes y descendentes se calcula sobre un total de 2600, a saber: $2000/2600$ y $600/2600$. Es decir: $0.77 + 0.23 = 1.00$.

Pero, en general, en análisis sociológicos, las proporciones son muy poco utilizadas porque los números decimales dificultan su lectura.

Porcentajes

Los porcentajes resultan de multiplicar por 100 el número de unidades en una posición del grupo experimental (o de la variable) y dividir el producto por el total de unidades de la variable (o del grupo experimental).

Para demostrar una de las ventajas de la utilización de los porcentajes, es conveniente volver al ejemplo inicial y presentarlo sólo en base a ellos.

CUADRO 11. Porcentaje de individuos según la dirección de su movilidad y su actitud hacia el partido político X

<i>Actitud</i>	<i>Movilidad descendente</i>	<i>Inmóviles</i>	<i>Movilidad ascendente</i>
Favorable	48	44	15
Indiferente	35	40	40
Desfavorable	17	16	45
Total	100% (N: 600)	100% (N: 6 400)	100% (N: 2 000)

Con respecto a este cuadro, debe notarse:

1) La exclusión de los números absolutos aclara notablemente la presentación y facilita la lectura, sin que por ello se pierda información, ya que el hecho de consignar los totales debajo de los 100% permite reconstruir el cuadro original calculando los números absolutos de cada posición con una simple operación inversa a la que llevó a obtener los porcentajes.

2) Los distintos grupos experimentales son comparables porque todos ellos han sido igualados (normalizados, estandarizados) en 100.

3) Aunque es difícil dar una regla de validez absoluta, los autores coinciden en que no deben calcularse porcentajes, a menos que el total de unidades sea de alrededor de 50 casos. Algunos autores dan una cifra de 20 a 30 casos. Pero todas estas cifras dependen del número de posiciones que haya en la variable (o en el grupo experimental). Puede ser válido calcular porcentajes de una dicotomía cuando el total de casos es 30, y no serlo calcular porcentajes de una variable con cinco posiciones cuando el total es 50.

4) En general, no deben incluirse los decimales, ya que con ello se pierde una ventaja del uso de los porcentajes (la facilidad y claridad de lectura). Sin embargo, los decimales deben calcularse cuando: a) el número de casos es muy grande y se justifica obtener una diferencia porcentual muy pequeña; b) es necesaria la precisión como sucede en la determinación de personas que han sanado (o enfermado) con la utilización de cierto medicamento; c) se prevé la realización de futuras mediciones, cuyos resultados no se conocen al practicar la medición actual. En este caso, hay que tratar de conservar las diferencias mínimas.

5) Por otra parte, los porcentajes son un instrumento que, si bien es superior a las razones y proporciones, es siempre elemental para el análisis, ya que sólo indican la dirección de la relación entre las variables, pero no la fuerza de la asociación entre ellas.

Cabe ahora preguntarse ¿cómo se calculan los porcentajes?, ¿cómo se leen los porcentajes?, ¿hay que volver a la lógica subyacente del diseño experimental, y a la propiedad-espacio?

El encabezamiento de cada columna indica cómo se han constituido los grupos experimentales. En el ejemplo, los grupos difieren en movilidad y en dirección de la

movilidad. Ésa es la variable independiente, aquella que para este estudio se considera que causa, determina, influye en la variable dependiente (actitud hacia el partido político X). La hipótesis adelanta que aquello en lo cual los grupos se diferencian (la dirección de la movilidad) es causa determinante o influyente en la variación o no variación de las respuestas al mismo estímulo, sea éste un indicador simple (una pregunta) o compuesto (un índice, una escala).

Teniendo en cuenta las consideraciones precedentes, es claro que los porcentajes deben calcularse en el sentido en que se han constituido los grupos experimentales (ya sea éste el sentido de las columnas o el de las hileras). Así se obtendrá la respuesta diferencial (y comparable) de cada grupo al estímulo. Se acostumbra a calcular los porcentajes en sentido vertical, porque se acostumbra a encabezar la columna con las variables consideradas independientes.

Por consiguiente, los porcentajes se leen en el sentido transversal a aquel en que se han calculado. Así se pueden ver las diferencias o las similitudes en frecuencias relativas de cada grupo experimental a los mismos grados del estímulo común.

Volviendo al ejemplo, se notará que la variable considerada independiente es la dirección de la movilidad. En base a ella se han constituido los grupos experimentales de individuos con movilidad descendente, los individuos inmóviles y los individuos con movilidad ascendente. Por eso los porcentajes se han calculado siguiendo las columnas. Pero la lectura procede en orden transversal a aquel en que se han calculado los porcentajes. Así se ve que los individuos con actitud favorable son un 48% de los descendentes, un 44% de los inmóviles y un 15% de los ascendentes. Por lo cual se puede concluir que los individuos con movilidad descendente tienen una actitud más favorable al partido político X que los individuos con movilidad ascendente. El mismo procedimiento puede seguirse con las hileras correspondientes a los indiferentes y a los que tienen actitud desfavorable.

Pero en realidad, dada la hipótesis que se había pedido comprobar, el cuadro se presenta con más datos que los necesarios. Tanto el cuadro que presenta los datos como el diseño de la propiedad-espacio deben ser lo más simples posibles con relación a la hipótesis. Como la hipótesis relacionaba sólo movilidad ascendente y descendente con actitud favorable y desfavorable, el cuadro inicial debería expresarse de la siguiente manera:

CUADRO 12

<i>Actitud hacia el partido político X</i>	<i>Movilidad descendente</i>	<i>Movilidad ascendente</i>
Favorable	75%	25%
Desfavorable	25%	75%
Total	100% (N: 390)	100% (N: 1200)

Pero como se ha señalado al considerar la propiedad-espacio, es muy poco lo que dice un análisis de dos variables. Si se agrega una tercera variable, el análisis se enriquece muchísimo.

Supóngase la hipótesis: “Cualquiera que sea el grado de desarrollo del país latinoamericano, vale la hipótesis de que los individuos móviles descendentes tienden a tener una actitud favorable hacia el partido político X y viceversa”:

CUADRO 13

<i>Grado de desarrollo</i>		<i>Bajo</i>			<i>Alto</i>		
<i>Actitud hacia el partido político X.</i>	<i>Dirección de la movilidad</i>	<i>Descendente</i>	<i>Ascendente</i>	<i>Total</i>	<i>Descendente</i>	<i>Ascendente</i>	<i>Total</i>
Favorable		100%	37%	48%	45%	10%	15%
		80	150	230	40	50	90
Desfavorable		—	63%	52%	55%	90%	85%
		—	250	250	50	450	500
		100%	100%	100%	100%	100%	100%
		80	400	480	90	500	590

Lo primero que conviene notar es que, si se pusiera en el título uno de los valores de la variable actitud hacia el partido político X, no se perdería información y el cuadro podría presentarse de la siguiente manera:

CUADRO 14. *Porcentaje de individuos con actitud favorable hacia el partido político X, según grado de desarrollo del país y dirección de la movilidad*

<i>Movilidad</i>	<i>Desarrollo bajo</i>	<i>Desarrollo alto</i>
Descendente	100	45
Ascendente	37	10

El porcentaje de individuos con actitud desfavorable hacia el partido político X resulta de la simple operación de restar 100 de los valores correspondientes a cada posición (0 en la celda *a*, 55 en la celda *b*, 63 en la celda *c* y 90 en la celda *d*). Nótese, al pasar, que los individuos con movilidad descendente de los países con bajo desarrollo aparecen teniendo más actitud favorable hacia el partido político X que los móviles descendentes de los países de alto desarrollo, y que los móviles ascendentes de países con bajo desarrollo también tienen una actitud más favorable hacia el partido político X que los móviles ascendentes de los países con alto desarrollo. Esto modifica las observaciones anteriores, y así aparece modificada la relación entre movilidad y actitud hacia el partido X por la introducción de la tercera variable: grado de desarrollo del país. Con esta observación podría aceptarse la hipótesis planteada, aunque modificada.

Conviene considerar más de cerca lo que sucede con la introducción de una tercera variable. Llamemos a la variable considerada más independiente (o condicionante, o antecedente) variable *x*; a la variable interpuesta, variable *t*, y a la variable dependiente, variable *y*. En este caso se trataría de grado de desarrollo del país, dirección de la movilidad y actitud hacia el partido político X, respectivamente.

GRÁFICA 4

X	X ₁			X ₂		
$\begin{array}{c} y \\ t \end{array}$	t ₁	t ₂	Total	t ₁	t ₂	Total
y ₁						
y ₂	3	4	No	3	4	No
Total	2	1		2	1	

Tal como hace Zeigel,⁷ el cuadro anterior podría esquematizarse a los fines de su presentación y análisis como aparece en la gráfica de la página anterior.

A partir de este esquema, conviene seguir los siguientes pasos para su presentación y análisis:

1) Presentación y análisis de los marginales de x_1 y x_2 . Según ejemplo, se trata de dos grupos de individuos móviles: uno con 480 individuos móviles de un país de bajo

desarrollo y otro con 590 individuos móviles de un país de alto desarrollo. Éstas son las celdas que en el esquema se señalan con el número 1.

2) Presentación y análisis de los marginales de la variable *t* (celdas número 2 del esquema), lo que origina el siguiente cuadro:

CUADRO 15

	<i>Bajo desarrollo</i>	<i>Alto desarrollo</i>
Movilidad descendente	17%	15%
Movilidad ascendente	83%	85%
	100%	100%
	(N: 480)	(N: 590)

De este cuadro se puede concluir que la dirección de la movilidad es similar en ambos tipos de países.

3) Presentación y análisis de las respuestas al estímulo de los distintos grupos experimentales: *a)* Presentación y análisis de las respuestas dadas por el grupo de movilidad descendente de países con bajo desarrollo y el grupo de movilidad descendente de países con alto desarrollo (celdas número 3 del esquema), lo que proporciona el siguiente cuadro:

CUADRO 16. Porcentaje de individuos con movilidad descendente, según el grado de desarrollo del país

<i>Actitud hacia el partido político X</i>	<i>Bajo desarrollo</i>	<i>Alto desarrollo</i>
Favorable	100%	45%
Desfavorable	—	55%
	100%	100%
	(N: 80)	(N: 90)

Del análisis de este cuadro puede concluirse que el grado de desarrollo del país tiene incidencia en la actitud hacia el partido político X entre individuos de movilidad descendente, ya que tienen actitud favorable 100% de los descendentes móviles de países con bajo desarrollo y sólo 45% de móviles descendentes de países con alto desarrollo, mientras que entre los móviles descendentes de países con bajo desarrollo no se registra actitud desfavorable y entre los móviles descendentes de países con alto desarrollo se registra 55% con actitud desfavorable. *b)* Presentación y análisis de las respuestas dadas por el grupo de movilidad ascendente de países con bajo desarrollo (celdas número 4 del esquema), lo que proporciona el siguiente cuadro:

CUADRO 17. *Porcentaje de individuos con movilidad ascendente, según el grado de desarrollo del país*

<i>Actitud hacia el partido político X</i>	<i>Bajo desarrollo</i>	<i>Alto desarrollo</i>	<i>Diferencia porcentual</i>
Favorable	87%	10%	+ 27
Desfavorable	63%	90%	- 27
	100% (N: 400)	100% (N: 500)	0

Del análisis de este cuadro puede concluirse que, para los móviles ascendentes, el grado de desarrollo del país también tiene incidencia en la actitud hacia el partido político X y que dicha incidencia se manifiesta en la misma dirección que para los móviles descendentes, aunque mucho más debilitada (obsérvese el comportamiento de la columna llamada diferencia porcentual).

Luego de estos análisis elementales se está en condiciones de decidir sobre la hipótesis planteada. Los datos permiten aceptar la hipótesis en el sentido de que, cualquiera que sea el grado de desarrollo del país, los individuos móviles descendentes tienden a tener una actitud favorable hacia el partido político X y viceversa. Pero con la importante especificación de que, si se mantiene constante la dirección de la movilidad, la hipótesis aparece especificada en el sentido de que la relación entre grado de desarrollo del país y actitud hacia el partido político X se acentúa entre los individuos móviles descendentes y se debilita entre los individuos móviles ascendentes.

Se observará que en el esquema hay celdas en las que se lee “No”. Tal vez sea innecesario decir que “No” no significa “Nunca”, pero conviene tenerlo presente. Esto quiere decir que, si en la constitución de los grupos experimentales se ha conservado o se estima que se ha logrado conservar la proporción que tales grupos tienen en el universo, entonces (y sólo entonces) es válido y lícito analizar las celdas correspondientes a los marginales de la variable dependiente (y) respecto de cada uno de los valores de la variable independiente (x). En el ejemplo, sólo si se supiera que la proporción de unidades medidas según los grados de desarrollo de los países es equivalente a la proporción de unidades en el universo de esos países, podría pensarse en analizar y_1 e y_2 para x_1 y x_2 . Como este dato se desconoce, no puede procederse a tal análisis. En todo caso, la presentación y análisis serían similares a los presentados anteriormente.

Medición del cambio porcentual

En numerosas ocasiones es necesario determinar el incremento porcentual que ha sufrido un valor de una variable en un grupo dado. Si se piensa en un programa o proyecto, deben instrumentarse modos de evaluar dicho programa o proyecto. Por ejemplo, se desea determinar la eficacia de un programa para difundir entre las madres el

conocimiento de ciertos elementos de higiene infantil. La mejor manera es constituir dos grupos de madres que sean lo más parecidos posibles en características básicas: nivel económico-social, nivel educacional, vivienda, disponibilidades diversas, etc. Uno de los grupos (el experimental) será sometido al programa, mientras que el otro grupo (el de control) no será sometido al programa. Al tiempo de iniciarse el experimento se mide el conocimiento que ambos grupos tienen de los elementos de higiene infantil que se quiere difundir. Esto proporciona un porcentaje (p_1) de madres que, antes del experimento, conocen los elementos de higiene. Después de sometido el grupo experimental al programa proyectado, se procede a una nueva medición que nos proporciona otro porcentaje de madres (p_2) que conocen los elementos de higiene.

El caso podría presentarse de la siguiente manera:

CUADRO 18

	<i>Grupo experimental</i>	<i>Grupo de control</i>
Nivel de conocimiento inicial (P_1)	10%	80%
Sometimiento al programa	sí	no
Nivel de conocimiento final (P_2)	20%	90%

Una forma de medir el cambio porcentual es calcular la *diferencia porcentual*, que se obtiene restando del nivel final el nivel inicial: $P_2 - P_1$. En el caso del ejemplo, se obtendría para el grupo experimental (de ahora en adelante GE) un incremento de 10% y para el grupo de control (de ahora en adelante GC), un incremento del 10%.

Otra manera de medir el cambio porcentual es mediante la aplicación de la *tasa de cambio*, es decir, dividiendo por el nivel inicial, la diferencia entre el nivel final y el nivel inicial: $\frac{P_2 - P_1}{P_1}$. En el caso del ejemplo, se tendría que el GE ha tenido un incremento porcentual de 100%, mientras que el GC sólo ha aumentado en 12.5%.

Pero, tanto la diferencia porcentual como la tasa de cambio porcentual tienen un defecto básico, que es prescindir del incremento potencial ($100 - P_1$). En el ejemplo, el GE parte de 10% y su incremento potencial es de 90%, mientras que el GC parte de 80% y su incremento potencial es de sólo 20%. Resulta claro que, para obtener una medida comparable del cambio porcentual, hay que incluir el incremento potencial en dicha medida.

A tal efecto, Hovland *et al.*⁸ exponen un *índice de efectividad* que resulta de dividir por el incremento potencial la diferencia entre el nivel final y el nivel inicial: $\frac{P_2 - P_1}{100 - P_1}$. En el ejemplo, medido por el índice de efectividad, el GE ha aumentado en 11% y el GC en 50%.

Las conclusiones a que conduce la aplicación de cada una de estas formas de medir el cambio porcentual son muy distintas y debe considerarse en cada caso cuál de ellas

corresponde aplicar.

Conviene analizar otro aspecto del índice de efectividad. Cuando se efectúan las mediciones en realidad no se encuentra sólo dos categorías de individuos: los que tienen la propiedad (en el ejemplo, conocer los elementos de higiene) y los que no tienen la propiedad, sino tres categorías de individuos, a saber: 1) Los que verdaderamente tienen la propiedad, a los que se denominará con la letra K (K_1 en la medición inicial, K_2 en la medición final). 2) Los que verdaderamente no tienen la propiedad, $100 - K_1$ en la medición inicial y $100 - K_2$ en la medición final. 3) Los que verdaderamente no tienen la propiedad, pero que por distintas razones se registran como teniéndola. Los individuos de esta categoría son una cierta proporción porcentual de los individuos de la categoría anterior: $x(100 - K_1)$ en la medición inicial y $x(100 - K_2)$, en la medición final.

Teniendo en cuenta estas categorías, los individuos que aparecen teniendo la propiedad en la medición inicial son: $K_1 + x(100 - K_1)$; y en la medición final son $K_2 + x(100 - K_2)$. Es decir,

$$P_1 = K_1 + x(100 - K_1)$$

y

$$P_2 = K_2 + x(100 - K_2)$$

Si se rempazan las nuevas expresiones para P_1 y P_2 en la fórmula del índice de efectividad, se obtiene:

$$\frac{P_2 - P_1}{100 - P_1} = \frac{K_2 + x(100 - K_2) - K_1 + x(100 - K_1)}{100 - K_1 + x(100 - K_1)}$$

$$1) \text{ Por asociatividad} = \frac{K_2 + x(100 - K_2) - K_1 - x(100 - K_1)}{100 - K_1 - x(100 - K_1)}$$

$$2) \text{ Por distributividad} = \frac{K_2 + 100x - xK_2 - K_1 - 100x - xK_1}{100 - K_1 - 100x - xK_1}$$

$$3) \text{ Por inverso} = \frac{K_2 - xK_2 - K_1 - xK_1}{100 - K_1 - 100x - xK_1}$$

$$4) \text{ Por distributividad} = \frac{K_2(1 - x) - K_1(1 - x)}{100(1 - x) - K_1(1 - x)}$$

$$5) \text{ Por cancelatividad} = \frac{K_2 - K_1}{100 - K_1}$$

Con esto, el índice de efectividad queda expresado en función de los individuos que verdaderamente tienen la propiedad. Obviamente, la aplicación de esta fórmula supone que, de alguna manera, se ha podido estimar la proporción de individuos que aparecen teniendo la propiedad, aunque en realidad no la tienen.

Los que se han planteado hasta el presente son los problemas más importantes del uso, cálculo, lectura y significación de los porcentajes. Queda todavía un problema menor. Es frecuente encontrar en las alternativas de respuesta para las preguntas de un cuestionario las alternativas *No sabe*, *No responde*, *No sabe/No responde*.

Estas alternativas pueden tener un doble significado: 1) Cuando se quiere averiguar si la unidad sabe o no sabe alguna cosa, la alternativa *No sabe* tiene un significado autónomo e importante para la investigación. 2) Cuando no se quiere averiguar si la unidad sabe o no sabe alguna cosa, las alternativas *No sabe*, *No responde*, *No sabe / No responde* son sólo categorías residuales introducidas a efecto de mantener la exhaustividad de la clasificación.

En el caso de tratarse de alternativas autónomas, lo normal es tratar de detectar cuáles son los motivos por los cuales las unidades no saben acerca del asunto en cuestión: porque no están informadas, porque están insuficientemente informadas, porque han recibido información distorsionada, o porque están evadiendo la respuesta a la pregunta formulada.

En el caso de tratarse de alternativas residuales pueden presentarse dos situaciones principales:

1) Las unidades que responden *No sabe* son pocas. En este caso puede prescindirse de la categoría.

2) Las unidades que responden *No sabe* son muchas. Las medidas previas que deben adoptarse para que las frecuencias de *No sabe* no sean elevadas hacen referencia a las instrucciones, a los entrevistadores y a analizar detenidamente los resultados de la prueba previa del cuestionario. Pero si todavía la frecuencia de individuos que responden *No sabe* es muy alta, pueden seguirse tres procedimientos que deben justificarse en cada caso: a) suponer que las respuestas *No sabe* se distribuyen al azar entre las restantes alternativas. Entonces se adjudica en forma proporcional a cada una de las alternativas restantes el total de individuos que respondieron *No sabe*. Este procedimiento es poco aconsejable; b) excluir directamente las respuestas *No sabe* y calcular los porcentajes para las alternativas restantes en base al nuevo total que resulta de excluir a los que respondieron *No sabe*; c) suponer que la respuesta *No sabe* es, en realidad, una evasiva. En este caso hay que buscar en el cuestionario alguna otra pregunta altamente asociada con aquella en que se han registrado numerosos *No sabe* y adjudicar los individuos que respondieron *No sabe*, según la respuesta dada en esa pregunta altamente asociada. Este procedimiento, que se podría llamar de consistencia de respuestas, es más aconsejable que los otros dos siempre que la asociación entre las respuestas a las preguntas sea realmente alta y que, en lo posible, se analice la asociación con más de una pregunta.

Bibliografía

- Hovland, C. J., *et al.*, “A Baseline for the Measurement of Percentage Change”, en Lazarsfeld, P., y Rosenberg, M., *The Language of Social Research*; The Free Press, 5ª ed., Nueva York, 1965.
- Galtung, J., *Teoría y métodos de la investigación social*; Eudeba, Buenos Aires, 1966.
- Barton, A. H., “The Concept of Property-Space in Social Research”, en Lazarsfeld, P., y Rosenberg, M., *The Language of Social Research*; The Free Press, Nueva York, 1955.
- Zeisel, H.: *Dígalos con números*; Fondo de Cultura Económica, México, 1961.

IX. ANÁLISIS DE DATOS: PAQUETE ESTADÍSTICO PARA LAS CIENCIAS SOCIALES (SPSS): OFERTA Y CONDICIONES PARA SU UTILIZACIÓN E INTERPRETACIÓN DE RESULTADOS

JORGE PADUA

EL ENORME desarrollo experimentado en los aparatos electrónicos en el almacenaje y procesamiento de datos, unido al acceso relativamente fácil de los investigadores a estas computadoras, abre para el científico social nuevas posibilidades tanto para el tratamiento de datos en gran escala cuanto para la aplicación de técnicas estadísticas complejas, en lo matemático o en lo referente a la cantidad de tiempo requerido para su cómputo manual.

El uso de computadoras para el análisis de datos requiere una serie de pasos:

a) Los datos deben ser preparados de manera tal que puedan ser “alimentados” a la máquina, es decir, los datos deben estar en la forma de tarjetas perforadas, cinta magnética, disco, o alguna otra forma especial (cintas de papel, por ejemplo).

b) El investigador debe decidir qué es lo que quiere hacer con sus datos, es decir, qué tipo de cálculo va a solicitar; correlaciones, coeficientes de confiabilidad, análisis de varianza, etcétera.

c) Es necesaria la preparación de un “programa”, que no es otra cosa que una guía de instrucciones que describe para la máquina la forma como los cálculos deben ser realizados.

d) Es necesario confeccionar además una serie de instrucciones para poner un conjunto particular de datos en un determinado centro de computación, para un determinado programa de cálculo. Esto es lo que se llama en la jerga técnica “*job*”. Típicamente cada *job* consiste de varios mazos de tarjetas: i) tarjetas de sistema (*system cards*); ii) programa; iii) control de tarjetas paramétricas, y iv) datos. Las *tarjetas de sistema* varían de un centro de computación a otro, y por lo general éstos instruyen a sus usuarios sobre la preparación de tarjetas, o se encargan ellos mismos de confeccionarlas. Estas tarjetas incluyen nombre del usuario, de la investigación, nombre de variables, etc. En la mayoría de los casos incluye además una tarjeta con el nombre del programa. El *programa* es la lista de instrucciones que especifican el tipo y el orden de las operaciones que la computadora va a efectuar. Cuando el programa está ubicado en la memoria de la computadora, se ubica por medio de una de las tarjetas del sistema. Las *tarjetas paramétricas* (o tarjetas de control) son una lista de instrucciones específicas para un juego de datos en particular y para el problema específico. Por lo general, estas tarjetas incluyen tarjetas de problemas, tarjetas opcionales y tarjetas de formatos. Las tarjetas de problemas contienen una descripción del *job* específico (número de

observaciones, número de variables, forma del *input* y del *output*, etc.). Las tarjetas de formatos instruyen a la máquina sobre dónde y cómo encontrar los datos en cada tarjeta, cinta o disco. Las tarjetas opcionales programan la máquina para hacer exactamente lo que debe hacer y con qué exactitud.

El objetivo específico de este trabajo es delinear algunos principios estadísticos y metodológicos referidos a los programas, particularmente a aquellos incluidos en el SPSS (*Statistical Package for the Social Sciences*), “Paquete estadístico para las ciencias sociales”, uno de los programas más completos y de mayor difusión en el medio académico y de investigación en el hemisferio occidental.

En la medida en que fueron los estadísticos y los psicometristas los que —de los científicos sociales— han utilizado más las computadoras, no es de extrañar que los programas más refinados se concentren en cálculos estadísticos; sin embargo, hay una generalización cada vez más creciente del uso de computadoras en una variedad de disciplinas (medicina, música, administración, simulación de relaciones internacionales, estrategia, etc.) que son una indicación de la flexibilidad de estas máquinas para realizar tareas complejas, al mismo tiempo que nos alertan sobre la necesidad de tener acceso a programadores y analistas de sistemas con cierto grado de familiaridad con los problemas sustantivos de disciplinas en particular y que sean capaces de preparar las computadoras para cálculos y operaciones no complementados en programas de tipo “paquete”.

Finalmente, quiero alertar al lector sobre las descripciones que aparecen en las secciones siguientes: me he preocupado de hacer accesible cada una de las posibilidades que aparecen en el SPSS; las más de las veces he seguido de cerca el estilo y el ejemplo de los autores del *Manual*, otras veces incluyo algo más de información. Todo el razonamiento es de carácter verbal, más que matemático, y la idea es presentar al investigador no familiarizado con la estadística matemática las opciones que hay en el programa, cuándo utilizarlas, cómo interpretar los *outputs*, etc. Para más detalles, en cada una de las diferentes técnicas existe abundante bibliografía que puede ayudar al lector a un uso más preciso del rico material de cálculo disponible en el SPSS.

La utilidad que puede tener este tipo de enfoque está en relación directa con el desconocimiento que el lector tenga de las limitaciones de cada uno de los estadísticos. Con demasiada frecuencia se solicita a los centros de computación cálculos que utilizan valioso tiempo de programación y de computadora para resultados finales sin ninguna significación. Por ejemplo, ocurre que el usuario solicita cálculos de correlaciones Pearsons para variables y atributos tales como sexo, opiniones, pertenencia a clase, etc. Ahora bien, una de las limitaciones de la correlación Pearsons es que las variables tienen que estar medidas a nivel intervalar al menos, y algunas de las variables que el usuario empleaba estaban medidas a nivel nominal u ordinal. ¿Quiere decir que no es posible utilizar correlaciones? De ninguna manera; simplemente lo que quiere decir es que había que computar el coeficiente de correlación o de asociación apropiado. Más adelante se analizarán algunas alternativas.

Otro problema bastante común y asociado a los niveles de medición es el de

seleccionar las estadísticas apropiadas para un conjunto determinado de datos, a partir del *output* de la computadora. Muchos programas (por ejemplo el subprograma *crosstab*) producen una serie de coeficientes de asociación, de los cuales el investigador debe seleccionar los que corresponden a sus datos; muchas veces ocurre que se publican todos o se utilizan algunos indebidamente.

En este trabajo, vamos a ocuparnos del paso *b* y en parte del paso *c* en el uso de computadoras. Es decir, vamos a proponer algunos estadísticos —la mayoría de los cuales están contenidos en un programa denominado SPSS— de manera tal que el investigador decida más apropiadamente qué quiere hacer con sus datos, qué tipo de cálculo va a solicitar, etcétera.

NIVELES DE MEDICIÓN

Uno de los requisitos teóricos más importantes para la utilización eficiente de modelos matemáticos o estadísticos es que éstos sean isomórficos con el concepto o el conjunto de conceptos que los modelos representan. En otras palabras, el modelo matemático debe tener la misma forma que el concepto. De no ser así, cualquier tipo de operación es ilegítima.

Las reglas para la asignación de números a objetos, conceptos o hechos están determinadas por una serie de teorías distintas, donde cada una de ellas se denomina nivel de medición. La teoría de la medición especifica las condiciones en que una serie determinada de datos se adaptan legítimamente a un nivel u otro, de manera que exista isomorfismo entre las propiedades de las series numéricas y las propiedades del objeto. De esta manera es posible utilizar el sistema matemático formal como un modelo para la representación del mundo empírico o conceptual.

Toda medición tiene tres postulados básicos, que son necesarios para igualar, ordenar y añadir objetos. Estos principios o postulados son: 1) $a = b$ o $a \neq b$, pero no ambos al mismo tiempo. 2) Si $a = b$ y $b = c$, entonces $a = c$. 3) Si $a > b$ y $b > c$ entonces $a > c$.

El primer postulado es necesario para la clasificación. Nos va a permitir determinar si un objeto es idéntico o no a otro en virtud del atributo que consideramos. Manteniendo constante la dimensión tiempo, establece relaciones excluyentes. El segundo postulado nos capacita para establecer la igualdad de un conjunto de elementos con respecto a una característica determinada. Es el principio de la transistividad de igualdades. El tercer principio, o principio de la transistividad de desigualdades o inecuaciones, nos permite establecer proposiciones ordinales o de rango.

Con base en estos postulados y de acuerdo con el tipo de operaciones empíricas que se puedan realizar con los atributos del universo que se desea escalar, hay cuatro tipos distintos de niveles de medición: *a*) Nominal, *b*) Ordinal, *c*) Intervalar, y *d*) Por cocientes o racionales.

Cada uno de estos niveles se caracteriza por el grado en que permanecen invariantes. La naturaleza de esta invariancia fija los límites a los modos de manipulación estadística que pueden aplicarse legítimamente a los datos incluidos en el nivel de medición.

A) Nivel nominal

Es la forma más elemental de medición, en la que simplemente se sustituyen los objetos reales por símbolos, números, nombres. Esta clasificación de los elementos de un universo de acuerdo a determinados atributos da a la medición a este nivel un significado más cualitativo que cuantitativo.

Para “medir” en este nivel, se asignan símbolos o signos al atributo del objeto o conjunto de objetos que se desea medir, con la condición básica de no asignar el mismo signo a categorías que son diferentes; o diferentes signos a la misma categoría.

Por medio de esta escala simplemente diferenciamos los objetos de acuerdo con la categoría a la que pertenecen. Ejemplos de medición de variable a nivel nominal:

Sexo (masculino-femenino).

Religión (católico, protestante, judío, mahometano, otra).

B) Nivel ordinal

En este nivel de medición, los objetos no solamente aparecen como diferentes, sino que además existe cierta relación entre grupos de objetos. Es decir, la relación “mayor que” es válida para todos los pares de objeto de diferente clase.

Se obtiene una escala ordinal “natural” cuando los datos originales admiten una relación “más grande que” para todos los pares de unidades. Los numerales asignados a los objetos rangueados son llamados valores de rango. Ejemplo: autoridad militar (capitán, teniente, sargento, etc.), distribución de poder o de prestigio, *status* socioeconómico (alto, medio, bajo).

C) Nivel intervalar

En las dos escalas examinadas más arriba, los elementos del sistema eran clases de objetos y las relaciones se reducían a igualdad o más grande que. Ninguno de los dos niveles especificaba *distancia* entre clases, es decir que cuando hablábamos de que A era mayor que B , y que B era mayor que C , no podíamos hacer ninguna afirmación sobre si la distancia que separaba A de B era mayor, igual, más o menos importante y cuán intensa, que la que separaba B de C . En lo que se refiere a distancia que separa objetos o clases de objetos tanto el nivel nominal como el ordinal son nominales.

En una escala intervalar podemos afirmar no solamente que tres objetos o clases a , b , c , están en una relación $a > b > c$, sino también que en los intervalos que separan los objetos se da la relación $\overline{ab} > \overline{ij}$ o $\overline{ij} > \overline{ab}$.

Es decir que es una escala o nivel que está caracterizado por un orden simple de los estímulos sobre la escala, y por un orden en los tamaños que miden las distancias en los estímulos adyacentes sobre la escala. Los datos contienen especificaciones relativas al tamaño exacto de los intervalos que separan a todos los objetos en la escala, además de las propiedades que se obtienen en la escala nominal y ordinal. Aquí estamos realmente a nivel de lo que entendemos por “cuantificación” propiamente tal y se requiere el

establecimiento de algún tipo de unidad física de medición que sirva como norma, y que por lo tanto pueda aplicarse indefinidamente con los mismos resultados.

La escala de intervalos supone la adjudicación de un cero arbitrario, y las operaciones aritméticas se aplican sobre las diferencias entre los valores de la escala. Ejemplos: temperatura, *tests* de *IQ*, etcétera.

D) Nivel por cociente o racional

Supone un 0 absoluto y es posible cuando existen operaciones para determinar cuatro tipos de relaciones: 1) Similitud. 2) Ordenación de rangos. 3) Igualdad de intervalos. 4) Igualdad de proporciones (razones o cocientes).

Una vez determinada la igualdad por cociente, los valores numéricos pueden transformarse con sólo multiplicar cada valor por una constante. Con este tipo de escala es posible realizar todo tipo de operaciones aritméticas. Ejemplo: distancia, peso, volumen, etcétera.

PROGRAMA ESTADÍSTICO DEL SPSS

El SPSS contiene programas estadísticos para: *a)* Estadística descriptiva y distribución de frecuencia para una variable. *b)* Tabla de contingencia y tabulaciones cruzadas. *c)* Correlaciones bivariatas. *d)* Correlaciones parciales. *e)* Regresiones múltiples. *f)* Análisis de la varianza. *g)* Análisis discriminatorio. *h)* Análisis factorial. *i)* Análisis de correlaciones canónicas. *j)* Análisis de escalograma, para escalas Guttman.

Y una serie de subrutinas, para modelos lineales en análisis de regresiones, como regresiones con variables mudas (*Dummy variables*) y *Path* análisis. Nosotros trataremos de especificar para qué sirven cada uno de estos subprogramas dando algunos detalles sobre las condiciones para su utilización.

ESTADÍSTICA DESCRIPTIVA

A) Medidas de tendencia central

Incluye solamente la media aritmética, la mediana y el modo. Estos *promedios* indican los valores centrales de observaciones. Sirven para: describir en forma sintética el conjunto de datos; los promedios provenientes de muestras pueden ser utilizados como una buena estimación de los valores parámetros, existiendo para ello una serie de técnicas de estimación a partir de valores muestrales que serán examinadas en la parte correspondiente a estadística inferencial.

La media aritmética es lo que conocemos familiarmente como promedio, esto es, el resultado de dividir la suma total de todas las mediciones por la cantidad total de casos. *La mediana* es el punto en la distribución que la divide en dos partes iguales, esto es, por encima de la mediana se encuentra el 50% de los casos y por debajo el otro 50%. *El modo* es el punto en la distribución que registra la frecuencia máxima. La media

aritmética es la más exacta y confiable de las tres medidas.

Empleo de media, mediana y modo

El nivel de medición apropiado para cada uno de los niveles es:

nominal	modo
ordinal	modo, mediana
intervalar o racional	modo, mediana, media

Existen algunos casos en los cuales, además del nivel de medición apropiado, es necesario tener en cuenta la forma de la distribución de los datos.

En síntesis, se computa la *media aritmética* cuando: *a)* Los datos están medidos a nivel intervalar al menos. *b)* Cuando la distribución es simétrica, aproximadamente normal y unimodal. *c)* Cuando se van a efectuar cálculos posteriores.

La *mediana* cuando: *a)* Los datos están medidos a nivel ordinal al menos. *b)* Cuando se cuenta con distribuciones incompletas. *c)* Cuando la distribución es necesariamente asimétrica.

El *modo* cuando: *a)* La escala es nominal. *b)* Cuando se desea conocer el caso más típico.

B) Medidas de variabilidad o de dispersión

Las medidas de tendencia central por sí solas constituyen una información valiosa, pero insuficiente para un análisis de la distribución, necesitando el complemento de lo que se conoce como medidas de variabilidad o de dispersión. Estas medidas indican cómo se distribuyen los valores alrededor de las medidas de tendencia central. Las medidas de variabilidad más importantes son:

Amplitud total (range), que denota simplemente la diferencia entre los valores máximo y mínimo de la distribución.

La amplitud semi-intercuartil (Q), que es la mitad de la amplitud de 50% central de casos.

La desviación media (AD) es la media aritmética de todos los desvíos con relación al promedio, cuando no se toman en consideración los signos algebraicos.

La desviación estándar (sigma, σ), que es un desvío cuadrático medio, o en términos operacionales la raíz cuadrada de la media aritmética del cuadrado de las desviaciones de cada una de las medidas en relación al promedio aritmético.

Cada una de estas medidas de variabilidad complementa la información de las medidas de tendencia central. Por ejemplo:

Medida de tendencia central

Medida de variabilidad

Modo

Mediana

Media

Amplitud total

Amplitud semiintercuartil

Desviación media-desviación estándar

Los valores provenientes de la *amplitud total* son útiles para tener una idea general del rango de variación en los datos; sin embargo, es poco confiable en la medida en que para su cálculo solamente se utilizan dos valores extremos, por lo cual es imprecisa.

La *amplitud semiintercuartil*, complemento de la mediana, es útil como índice de la simetría de la distribución total. En distribuciones perfectamente simétricas el cuartil 1 (Q_1) y el cuartil 3 (Q_3) están a igual distancia del centro de la distribución o mediana (Q_2). Si las distancias son desiguales hay asimetría. En resumen:

Asimetría positiva cuando: $(Q_3 - Q_2) > (Q_2 - Q_1)$

Asimetría negativa cuando: $(Q_3 - Q_2) < (Q_2 - Q_1)$

Asimetría cero cuando: $(Q_3 - Q_2) = (Q_2 - Q_1)$

La *desviación media* nos informa sobre la magnitud de las desviaciones respecto a la media. Cuando la distribución es normal, aproximadamente 58% de las observaciones caen en el espacio comprendido entre la media aritmética más una desviación media y la media aritmética menos una desviación media.

La *desviación estándar* es la medida de variabilidad más exacta y confiable y la más empleada en cálculos posteriores (correlación, varianza, etc.). La interpretación más común de la desviación estándar es idéntica a la que realizamos con la desviación media, esto es, en términos de distribución normal; sumando y restando a la media aritmética una desviación estándar debe esperarse que el 68.26% de todos los casos caiga en esa área de la curva. De esta manera podemos estimar para una distribución cualquiera cuánto se acerca o se aleja de una distribución normal. Cuando la distribución es muy asimétrica el cálculo de la desviación estándar no es conveniente, recomendándose más el cálculo de desviación media o desviación intercuartil.

El SPSS contiene además una serie de medidas para la determinación de la forma de la distribución como la *curtosis* (*kurtosis*) y la *oblicuidad* (*skewness*).

La *oblicuidad* es un estadístico que indica el grado en que la distribución se aproxima a la distribución normal. Cuando la distribución es completamente simétrica, la oblicuidad es igual a 0. Los valores positivos indican que los casos se concentran más a la izquierda de la media, mientras que los valores extremos, a su derecha. Los valores negativos se interpretan exactamente al revés. Se aplica únicamente cuando los datos están a nivel intervalar al menos.

La *curtosis* es una medida relativa a la forma de la distribución (mesocúrtica o platicúrtica). La curtosis en una distribución normal es cero. La curtosis es positiva cuando la distribución es estrecha y en forma de pico, mientras que los valores negativos

indican una curva aplanada.

El cálculo de estadística descriptiva en el SPSS está contenido en dos subprogramas: *condescriptive* y *frecuencias*. El subprograma *condescriptive* en variables continuas a nivel intervalar. El subprograma *frecuencias* trabaja con variables discretas, por consiguiente se corresponde a los niveles nominales y ordinales.

El investigador debe tener cuidado en la selección de los valores que aparecen en las descripciones sumarias en las hojas del *out-put* de la computadora, con el fin de seleccionar solamente aquellos estadísticos que se corresponden con la naturaleza de sus datos.

CONFIABILIDAD DE LOS ESTADÍSTICOS

Las características de las poblaciones (o universos) son denominados valores parámetros. Qué es una *población*, desde el punto de vista estadístico, es una materia de definición arbitraria, aunque en un sentido general puede definirse como el conjunto total de unidades (individuos, objetos, reacciones, etc.) que pueden ser delimitados claramente por la posesión de un atributo o cualidad propio y único. *Muestra* es cualquier subconjunto de ese conjunto mayor que constituye la población. Téngase bien claro entonces que dos muestras pueden considerarse como provenientes de la misma población o de dos poblaciones diferentes, en relación a un solo aspecto (por ejemplo, coeficiente intelectual, actitud frente al cambio, etcétera).

Gran parte de las investigaciones en ciencias sociales se basan en estudio de muestras, y a partir de ellas se desea estimar o inducir las características de la población o universo.

El *error estándar* nos da una estimación de la discrepancia de las estadísticas muestrales con relación a los valores de la población. Si pretendemos utilizar el estadístico muestral como una estimación de los valores parámetros, cualquier desviación de la media muestral sobre la media parámetro puede considerarse como un error.

El *error estándar de la media aritmética* nos dice cuál es la magnitud del error de estimación para la media aritmética. El error estándar de la media es el desvío estándar de la distribución de las medias muestrales provenientes de una misma población. El error estándar de la media es directamente proporcional al tamaño de su desviación estándar e indirectamente proporcional al tamaño de la muestra. La interpretación del error estándar de una media se realiza en términos de distribuciones normales; esto es, la finalidad es determinar en qué medida las medias aritméticas de muestras similares a la que estamos considerando se alejan o aproximan a la media de la población. Cuanto menor sea el valor del error estándar tanto mayor será la fiabilidad de nuestra media como estimación de la media de la población. En general, la probabilidad de error se fija con valores de .05 y .01, es decir que la probabilidad de error aceptada es de 5 sobre 100 o de 1 sobre 100, de ahí que la estimación de los parámetros se haga en *intervalos de confianza*, donde a la media muestral le agregó y le disminuyó tantos valores de error

estándar como para cubrir 95% del total de casos, o el 99%.

La fórmula genérica para la determinación de intervalos de confianza es:

Para	.05	:M	±	1,96	σ_M
Para	.01	:M	±	2,58	σ_M

Para la distribución de probabilidades distintas de .05 y .01, simplemente se buscan en la tabla de curva de distribución normal los valores correspondientes.

Las interpretaciones señaladas tanto para la media aritmética como para el resto de los estadísticos son posibles únicamente cuando las muestras son aleatorias.

El cálculo del error estándar de la media está contenido en el subprograma *condescriptive*.

Confiabilidad de la mediana. Se interpreta de la misma forma que el error estándar de la media, y en poblaciones normalmente distribuidas, el error estándar de la mediana da una variabilidad aproximadamente de 25% mayor que el de las medias.

Error estándar de desviación estándar, Q, proporciones, porcentajes, frecuencias. Las fórmulas para su cálculo difieren, pero la interpretación es idéntica a la realizada más arriba, esto es, en términos de estimación de valores parámetros por medio de intervalos de confianza.

Confiabilidad de un coeficiente de correlación. Lo mismo que los demás estadísticos, un coeficiente de correlación está sujeto a los errores de muestreo. La variabilidad entonces estará en función del tamaño del error estándar del coeficiente. Sin embargo, la distribución de los coeficientes obtenidos no será uniforme, ya que dependerá de la magnitud de r , así como el número de observaciones que componen la muestra. En la medida en que los coeficientes varían entre + 1.00 y — 1.00 cuando la r parámetro se aproxima a esos valores extremos, la distribución será más asimétrica; negativamente asimétrica para los valores positivos de r y positivamente asimétrica para los valores negativos de r . Solamente en el caso en que el r parámetro sea 0, entonces la distribución de las r muestrales será normal.

Para muestras grandes el problema de la asimetría carece de importancia significativa; cuanto más grande la muestra, menor será la dispersión de las r , de ahí que, aun cuando r es cero, en muestras menores de 25 casos sea necesario tener alguna precaución en las estimaciones. Cuando r es grande y la muestra es grande, el error estándar del coeficiente será mínimo.

CONFIABILIDAD DE DIFERENCIA ENTRE ESTADÍSTICOS (SUBPROGRAMA BREAKDOWN)

El subprograma *breakdown* calcula e imprime las sumas, medias, desviaciones estándar y varianza de la variable pendiente en los distintos subgrupos que la componen, según clasificaciones complejas que incluyan de 1 a 5 variables independientes, cualquiera que sea el nivel de medición (en las variables independientes, la variable dependiente debe

ser medida a nivel intervalar).

Antes de desarrollar las alternativas del subprograma de *breakdown* conviene introducir conceptualmente la idea de confiabilidad de la diferencia entre estadísticos.

Para la investigación es importante no solamente estimar los valores poblacionales, sino utilizar el error estándar para interpretar varios resultados, en lo relativo a las diferencias que pueden existir entre ellos.

El tipo de preguntas que nos planteamos aquí es: ¿Cuál es la fiabilidad de la diferencia entre medias proporciones, etc., que hemos registrado en nuestras observaciones? ¿Son los hombres o las mujeres más capaces en comprensión verbal? ¿El rendimiento intelectual en las clases medias es superior, inferior o igual al rendimiento intelectual en las clases bajas?, etcétera.

A) Error estándar de la diferencia de medias (Subprograma T-Test)

La magnitud de la oscilación en la diferencia entre medias obtenidas de muestras distintas dependerá naturalmente de la magnitud de la oscilación que es propia de las medias. La estabilidad de las medias estará representada por sus respectivos errores estándares.

Cuando las N son lo suficientemente grandes, las medias oscilan alrededor de un valor central (parámetro que por lo general no conocemos). Nuestra finalidad es entonces determinar primero si existe diferencia, para luego definir su magnitud.

El problema reside entonces en determinar si la diferencia que se examina entre las dos medias muestrales implica además una diferencia en la distribución de la población; en otras palabras, si la diferencia es la expresión de diferencias reales a niveles poblacionales, o se deben simplemente a los efectos del azar (y por consiguiente del error) en las muestras.

El *test T* de student nos ayuda a establecer cuándo la diferencia entre dos medias es significativa. Para ello se formula una hipótesis nula. Una hipótesis nula supone que las dos muestras provienen de la misma población; consecuentemente las desviaciones son interpretadas como debidas al efecto del azar sobre las muestras. Según la hipótesis nula se supone que la distribución de las diferencias es normal, de donde $M_1 - M_2 = 0$.

El nivel de significación para la aceptación o rechazo de la hipótesis nula es seleccionado por el investigador. Los más comunes son de .05 y .01 aunque esto depende más bien del área que se está investigando (para aceptar o rechazar una vacuna nueva que cure el cáncer puedo elegir un nivel menor; cuando se trata de la introducción de una medicina para suplantarlo alguna en uso con cierto grado de eficacia, elegiré un nivel de significación mayor).

El valor de t nos informará entonces sobre la probabilidad o improbabilidad para la aceptación o rechazo de la hipótesis nula, o de alguna hipótesis alternativa. Es decir, no se afirma que no existe una diferencia en los resultados, sino únicamente que la diferencia es, o no es, significativa.

El subprograma *T-Test* computa los valores t y sus niveles de probabilidad para dos

tipos de casos: a) *Muestras independientes* o error estándar de la diferencia para medias no correlacionadas, es decir, para situaciones en las que las dos series de observaciones son independientes. Por ejemplo, comparación del rendimiento de hombres y mujeres en una situación de *test*. b) *Muestras apareadas* o error estándar de la diferencia para medias correlacionadas. El ejemplo típico es el de las mediciones antes-después en diseños experimentales.

Existen casos en los cuales el investigador no plantea la hipótesis nula (la hipótesis de las no-diferencias), sino que plantea una hipótesis alternativa, en la que trata de demostrar que la media en un grupo es más grande que la media del otro. En estos casos la interpretación del valor *t* obtenido se hace a partir de lo que se llama *test* de una sola cola, es decir, se toma en cuenta solamente una mitad de la distribución.

TABLAS DE CONTINGENCIA Y MEDIDAS DE ASOCIACIÓN (SUBPROGRAMA CROSSTABS)

Este subprograma contiene tanto tabulaciones cruzadas para tablas de $n \times k$ como una serie de medidas de correlación, asociación y confiabilidad de la diferencia entre estadísticos. Comenzaremos por los análisis de tipo más sencillo para continuar luego con los cálculos de medidas más complejas.

Tabulaciones cruzadas

Una tabulación cruzada es simplemente la combinación de dos o más variables discretas o clasificatorias en la forma de tabla de distribuciones de frecuencia. El cuadro resultante puede ser sometido a análisis estadístico, en términos de distribuciones porcentuales, aplicación de *test* de significación, coeficientes de asociación y de correlación, etcétera.

Las tablas cruzadas son muy utilizadas en análisis de encuestas, en tablas de 2×2 o con la introducción de variables de prueba o de control o intervinientes, constituyendo así tablas de $n \times k$. Aquí el investigador debe tener especial cuidado cuando solicita tablas de $n \times k$ de que el tamaño de su muestra sea lo suficiente grande para permitir que cada uno de los casilleros contenga las frecuencias esperadas.

De todos modos existe una serie de restricciones al uso de los distintos estadísticos que señalaremos más adelante y que imponen algunas limitaciones en cuanto a la cantidad total de casilleros, ya sea por cantidad de variables o por cortes en cada una de ellas. Por ejemplo: el investigador debe recordar que, si combina, digamos, 4 variables, todas dicotomizadas, la cantidad total de casilleros será de 16; si las variables estuvieran tricotomizadas la cantidad de casilleros ascendería a 81. La fórmula genérica para el cálculo del tamaño final de la matriz es:

$$M = r_1 \cdot r_2 \cdot r_3 \dots r_n$$

donde:

r = Cantidad de cortes o divisiones en cada una de las variables
 M = Tamaño de la matriz de datos

Es de recordar que en este tipo de cuadros hay que esperar un promedio de 10 a 20 casos en cada uno de los casilleros, lo que hace que las muestras deben tener tamaños considerables cuando se desea cuadros muy complejos.

Normalmente los cuadros imprimen tanto las frecuencias dentro de cada casillero o celda como los porcentajes con respecto al marginal horizontal y al marginal vertical y al total general (en ese orden), además de todos los coeficientes que incluye la subrutina y que a continuación pasamos a detallar.

Ji cuadrado (χ^2)

Es un modelo matemático o *test* para el cálculo de la confiabilidad o significado de diferencias entre frecuencias esperadas (f_e) y frecuencias observadas (f_o). La utilidad de este *test* no-paramétrico para variables nominales reside en su aplicación para prueba de hipótesis para tres tipos de situaciones:

a) Prueba de hipótesis referidas al grado de discrepancia entre frecuencias observadas y frecuencias esperadas, cuando se trabaja sobre la base de principios apriorísticos.

b) Pruebas de hipótesis referidas a la ausencia de relación entre dos variables. Se trata de pruebas de independencia estadística y son trabajadas en base a cuadros de contingencia.

c) Pruebas referidas a la bondad de ajuste. En este caso se trata de comprobar si es razonable aceptar que la distribución empírica dada (datos observados) se ajusta a una distribución teórica, por ejemplo, binomial, normal, Poisson, etc. (datos esperados).

Supuestos y requisitos generales

a) Las observaciones deben ser independientes entre sí. b) Los sucesos deben ser mutuamente excluyentes. c) Las probabilidades que figuran en las tablas de χ^2 están basadas en una distribución continua, mientras que el χ^2 calculado en la práctica lo está en base a variables discretas. Se supone que esta última puede aproximarse a la primera. d) El nivel de medición mínimo es nominal. e) Las frecuencias esperadas mínimas por casillero deben ser 5, cuando esto no se cumple es necesario aplicar un factor de corrección (corrección de Yates). f) La prueba de χ^2 es útil solamente para decidir cuándo las variables son independientes o relacionadas. No nos informa acerca de la intensidad de la relación, debido a que el tamaño de la muestra y el tamaño del cuadro ejercen una influencia muy fuerte sobre los valores del *test*. Existen numerosos estadísticos basados en la distribución de χ^2 que son útiles para la determinación de la intensidad de la relación (ver coeficiente ϕ , Cramer, C , etcétera).

Coeficiente ϕ (ϕ)

Es una medida de asociación (fuerza de la relación) para tablas de 2×2 . Toma el valor cero cuando no existe relación, y el valor + 1.00 cuando las variables están perfectamente relacionadas.

Coeficiente V. de Cramer

Es una versión ajustada del coeficiente ϕ para tablas de $r \times k$. El nivel de medición es nominal y el coeficiente varía entre 0 y 1.00.

Coeficiente de contingencia (C)

Basado como los dos anteriores en χ^2 se pueden utilizar matrices de cualquier tamaño. Tiene un valor mínimo de 0, y sus valores máximos varían según el tamaño de la matriz (por ejemplo, para matrices de 2×2 el valor máximo de C es .707; en tablas de 3×3 es .816, etc., la fórmula genérica: $\sqrt{\frac{k-1}{k}}$) consecuentemente para una interpretación del coeficiente obtenido en cualquier tabla de 2×2 habría que dividir ese valor entre .707.

Limitaciones: *a)* El límite superior del coeficiente está en función del número de categorías; *b)* dos o más coeficientes C no son comparables, a no ser que provengan de matrices de igual tamaño.

El coeficiente Q de Yule

También como los anteriores para escalas nominales, se utiliza únicamente en tablas de 2×2 . Los valores Q son 0 cuando hay independencia entre las variables, siendo sus límites ± 1.00 cuando cualquiera de las 4 celdas en el cuadro contiene 0 frecuencias: en general, cuál de los distintos coeficientes es preferible en este caso (ϕ o Q) depende del tipo de investigación y del tipo de distribución marginal.

Coeficiente lambda (λ)

Es un coeficiente de asociación para tablas de $r \times k$, cuando las dos variables están medidas a nivel nominal.

El coeficiente lambda pertenece a la familia de un grupo de coeficientes (τ_b , λ y otros), que se utilizan para hacer interpretaciones probabilísticas en cuadros de contingencia. El tamaño del coeficiente indica la reducción proporcional en errores de estimación en la variable dependiente cuando los valores en la variable independiente son conocidos.

El valor máximo de λ es 1.00 y ocurre cuando las predicciones pueden hacerse sin ningún error. Un valor cero significa que no hay posibilidad de mejorar la predicción. Un coeficiente lambda .50 significa que podemos reducir el número de errores a la mitad, etcétera.

Coefficiente τ_b de Goodman y Kruskal

Sirve a los mismos propósitos que el coeficiente lambda y debe preferirse cuando los marginales totales no son de la misma magnitud.

Coefficiente de incertidumbre

También para niveles nominales en cuadros de contingencia de $r \times k$. La computación del coeficiente toma en cuenta simetría y asimetría (el coeficiente lambda toma en cuenta, por ejemplo, solamente la asimetría). El coeficiente asimétrico es la proporción de reducción de la incertidumbre conocido por efecto del conocimiento de la variable independiente. La ventaja de este coeficiente sobre lambda es que considera el total de la distribución y no solamente el modo.

El máximo valor del coeficiente de incertidumbre es 1.00 que denota la eliminación de la incertidumbre, y se alcanza cada vez que cada categoría de la variable independiente está asociada solamente a una de las categorías de la variable dependiente. Cuando no es posible lograr ningún avance en términos de disminución de la incertidumbre, el valor del coeficiente es 0. Una versión simétrica del coeficiente mide la reducción proporcional en incertidumbre que se gana conociendo la distribución conjunta de casos.

Coefficiente tau b

Mide la asociación entre dos variables ordinales en cuadros de contingencia. Este coeficiente es apropiado para cuadros cuadrados (es decir, donde el número de columnas es idéntico al número de filas). Sus valores varían de 0 a ± 1.00 . El valor cero indica que no existe asociación entre pares concordantes y discordantes. El valor ± 1.00 se obtiene cuando todos los casos se ubican a lo largo de la diagonal mayor. En tablas de 2×2 el valor de tau b es idéntico al de ϕ con la ventaja de que el coeficiente tau b proporciona información sobre la dirección de la relación a través del signo. Los valores negativos indican que los casos se distribuyen sobre la diagonal menor. Los valores intermedios entre 0 y 1 indican casos que se desvían de las diagonales. A mayor desviación mayor proximidad al valor cero (es decir, cuando los pares discordantes son iguales a los pares concordantes).

Coefficiente tau c

Sirve a los mismos propósitos que el coeficiente tau b , pero este coeficiente es más apropiado para cuadros rectangulares (cuando el número de columnas difiere del número de líneas). La interpretación de ambos coeficientes es similar.

Coefficiente gamma (γ)

Mide asociación entre dos variables ordinales en cuadros de contingencia de $r \times k$.

Mientras que el coeficiente tau c depende para su cómputo solamente del número de líneas y de columnas, y no de las distribuciones marginales, tomando en cuenta los empates, el coeficiente gamma excluye los empates del denominador de la fórmula de cómputo, siendo además un coeficiente con posibilidades de aplicación en datos no agrupados. Además, el coeficiente no requiere cambios en la forma de la matriz. Los valores numéricos de gamma por lo general son más altos que los valores de tau b y de tau c .

El coeficiente gamma es simplemente el resultado del número de pares concordantes menos el número de pares discordantes, divididos por el número total de pares unidos. Los valores gamma varían entre 0 y ± 1.00 , donde el signo indica la dirección de la relación y los valores la intensidad de la misma.

El SPSS provee valores gamma para cuadros de tres a n entradas, en el que se calcula el gamma de orden cero y además gammas parciales. El gamma de orden cero mide la relación entre dos variables, siendo exactamente el mismo que se discute en los párrafos anteriores. Cuando la matriz tiene tres o más dimensiones, el SPSS (subprograma *crosstabs*) computa un coeficiente gamma de orden cero (reduciendo la tabla a variable dependiente e independiente) y, además, medidas de correlación parcial gamma de la relación entre las dos variables, controladas por una o más variables adicionales. El investigador puede analizar así cómo influye en la relación de sus variables dependiente e independiente la introducción de variables adicionales (en la sección correspondiente a correlaciones parciales indicaremos con mayor detalle el uso y significado de las correlaciones parciales).

Coeficiente D de Sommer

Para variables ordinales en cuadros de contingencia, este coeficiente toma en cuenta los empates, pero el ajuste es realizado de manera distinta de la utilizada en los coeficientes tau b y tau c .

Coeficiente eta (η)

Se utiliza cuando la variable independiente es nominal y la variable dependiente intervalar. Este coeficiente indica cuán disimilares son las medias aritméticas en la variable dependiente dentro de las categorías establecidas por la variable independiente. Cuando las medias son idénticas el valor del coeficiente es 0. Si las medias son muy diferentes y sus varianzas son pequeñas, los valores de eta se aproximan a 1.00.

Correlación biserial (r_b)

Para utilizar cuando una de las variables está medida a nivel nominal y la otra a nivel intervalar, la variable a nivel nominal puede ser una dicotomía forzada. Sus valores oscilan entre 0 y ± 1.00 .

Correlación punto-biserial (r_{pb})

Similar al coeficiente biserial, se aplica cuando la variable nominal es una dicotomía real. La interpretación de ambos coeficientes es idéntica y su utilización más común se encuentra en la construcción de pruebas, sobre todo para la determinación de validez.

Coeficiente de correlación Spearman (ρ)

Es un coeficiente de correlación por rangos, cuando las dos variables están medidas a nivel ordinal, e indica el grado en que la variación o cambio en los rangos de una de las variables está relacionado con las variaciones o cambios en los rangos en la otra variable. Tanto el coeficiente ρ (rho) de Spearman como el coeficiente τ (tau) de Kendall son coeficientes no paramétricos, es decir que no se hacen supuestos acerca de la distribución de los casos sobre las variables. Ambos coeficientes suponen la no existencia de muchos empates, por lo cual los sistemas de organización de los datos y de cómputo son distintos de los de las tabulaciones cruzadas, y por ello se encuentran en subprogramas diferentes (en este caso el subprograma correspondiente en el SPSS es denominado *nonpar corr*).

Para el cómputo del coeficiente correlación ρ de Spearman (así como para el τ de Kendall) no se toman en consideración los valores absolutos en las variables, sino su orden de rango. El coeficiente rho de Spearman se aproxima más que el coeficiente tau de Kendall al coeficiente de correlación producto-momento de Pearson, cuando los datos son aproximadamente continuos. Los valores del coeficiente varían entre -1.00 y $+1.00$.

Coeficiente de correlación tau de Kendall (τ)

Similar al coeficiente rho, se utiliza cuando las dos variantes son ordinales. Por lo general debe preferirse cuando existe abundante número de empates entre rangos, lo que se da especialmente cuando el número total de casos es grande y se clasifican en un número relativamente pequeño de categorías. El subprograma *nonpar corr* contiene factores de corrección para empates tanto para el coeficiente tau como para el coeficiente rho. Los valores de este coeficiente oscilan entre -1.00 y $+1.00$.

Coeficiente de correlación producto-momento de Pearson (r)

Para dos variables medidas a nivel intervalar por lo menos, éste es un coeficiente de correlación paramétrico que nos indica con la mayor precisión cuándo dos cosas están correlacionadas; es decir, hasta qué punto una variación en una se corresponde con una variación en otra. Sus valores varían de $+1.00$, que quiere decir correlación positiva perfecta, a través de 0 , que quiere decir independencia completa o ausencia de correlación, hasta -1.00 , que significa correlación perfecta negativa. El signo indica por lo tanto la dirección de la covariación y la cifra, la intensidad de la misma. Una

correlación perfecta de +1.00 indica que, cuando una variable se “mueve” en una dirección, la otra se mueve en la misma dirección y con la misma intensidad. La interpretación de la magnitud de r depende en buena medida del uso que se quiera dar del coeficiente, el grado de avance teórico en el área, etc. Guilford¹ sugiere como orientación general la siguiente interpretación descriptiva de los coeficientes de correlación producto-momento:

r menor que .20 — correlación leve, casi insignificante
 r de .20 a .40 — baja correlación, definida, pero baja
 r de .40 a .70 — correlación moderada, sustancial
 r de .70 a .90 — correlación marcada, alta
 r de .90 a 1.00 — correlación altísima, muy significativa

De todos modos, la interpretación del coeficiente está además condicionada a su grado de significación (ver significación de los estadísticos).

Premisas o suposiciones fundamentales para el cómputo de r : a) Ambas variables deben ser medidas a nivel intervalar al menos. b) La dirección de la relación debe ser rectilínea. c) La distribución tiene que ser homoscedástica (las dispersiones en las columnas y en las líneas del diagrama de dispersión deben ser similares). Esta condición prevalece cuando las dos distribuciones son simétricas entre ellas.

El programa imprime el valor del coeficiente de correlación, la cantidad de casos y la significación estadística.

Existen varios coeficientes que se derivan del coeficiente de correlación producto-momento, entre otros, por ejemplo: r^2 : mide la proporción de la varianza en una variable que es “explicada” por la otra.

Diagrama de dispersión (Scattergram)

El SPSS puede imprimir, además, a través de su subprograma *Scattergram*, el diagrama de dispersión para dos variables, computando además la regresión linear simple. El diagrama de dispersión es un tráfico de puntos donde, basado en los valores en las dos variables, una de las variables define el eje horizontal y la otra el eje vertical. Estos diagramas son de mucha utilidad, ya que nos dan una imagen de la relación, que puede ser utilizada para la determinación de la homocedasticidad, por ejemplo, y para decidir si vale o no la pena continuar más adelante.

Para la confección de los diagramas, el usuario tiene que tomar algunas decisiones sobre cómo va a manejar la falta de datos (*missing data*), qué clase de escala tiene que utilizar y cómo colocará las líneas segmentadas.

Comúnmente, dos líneas verticales y dos líneas horizontales segmentadas dividen cada eje con tres secciones, de manera que el gráfico consta de 9 rectángulos iguales. Si el investigador prefiere, las líneas segmentadas pueden ser diagonales que atraviesan el gráfico.

Los datos (es decir, cada punto sobre el diagrama) están representados por asteriscos (*) cuando un caso cae en alguna intersección, de dos a ocho casos el número es impreso. Nueve o más casos están representados por el número 9. Cuando la escala contiene muy pocas categorías, existe la posibilidad de que los puntos sobre el diagrama se den muy amontonados, lo que limita la utilización del diagrama de dispersión recomendándose para esas situaciones una tabulación cruzada.

Los estadísticos que acompañan al diagrama de dispersión son aquellos asociados a las regresiones lineares simples: correlación producto-momento, error estándar de la estimación, r^2 , significación de la correlación, intersección con el eje vertical, e inclinación.

Es necesario discutir con algún detalle el concepto de regresión, ya que sirve de base para la utilización de predicciones, así como de ayuda para la comprensión del concepto de correlaciones parciales y múltiples.

El concepto de regresión trata de describir no solamente el grado de relación entre dos variables, sino *la naturaleza* misma de la relación, de manera que podamos predecir una variable conociendo la otra (por ejemplo, el rendimiento académico a partir del resultado en un *test*, el ingreso a partir de la educación, etc.). Aquí no estamos interesados en explicar por qué las variables se relacionan como se relacionan, sino simplemente, a partir de la relación dada, predecir una variable a partir del conocimiento de los valores en la otra. Si la variable X es independiente de la variable Y (es decir, si son estadísticamente independientes), no estamos en condiciones de predecir Y a partir de X o viceversa, es decir, nuestro conocimiento de X no mejora nuestra predicción de Y . Por razonamiento inverso, cuando las variables son dependientes —están correlacionadas, co-varían—, el conocimiento de X nos puede ayudar a predecir el comportamiento de Y y viceversa.

Esto se logra mediante lo que se llama ecuación de regresión de Y sobre X , que nos da la forma en cómo las medias aritméticas de los valores de Y se distribuyen según valores dados de X .

La operación de regresión contiene los siguientes supuestos: *a)* Que la forma de la ecuación es lineal. *b)* Que la distribución de los valores de Y sobre cada valor de X es normal. *c)* Que las varianzas de las distribuciones de Y son similares para cada valor de X . *d)* Que el error es igual a 0.

Cumplidas estas condiciones, la ecuación de la regresión es:

$$Y = \alpha + \beta X$$

donde α y β son constantes y se les da una interpretación geométrica.

Si X es igual a 0, entonces $Y = \alpha$.

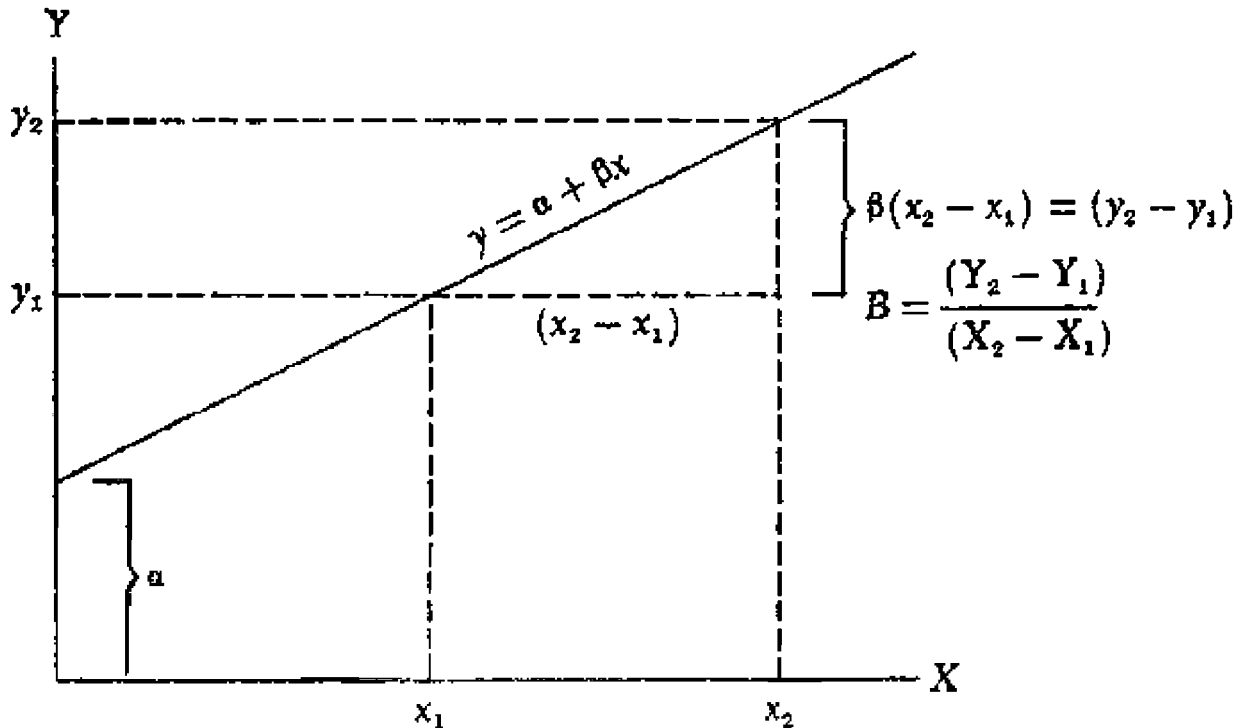
α representa entonces el punto donde la línea de regresión cruza el eje de Y .

La inclinación de la línea de regresión es dada por β , indicando la magnitud en el cambio de Y por cada unidad de cambio en X . Cuando β es igual a 1, y si las unidades de X y Y están indicadas por distancias idénticas a lo largo de sus ejes respectivos, la línea

de regresión estará en un ángulo de 45° con respecto al eje de las X . A más grande el tamaño de β , mayor será el declive, es decir, más grande el cambio en Y dados determinados valores de cambio en X .

H. Blalock ² presenta la siguiente figura que aclara la interpretación geométrica del coeficiente de regresión:

FIGURA 1



Es decir que β mide la tangente del ángulo, con lo cual queda identificado el ángulo.

Correlación parcial

Todos los coeficientes de correlación y asociación examinados hasta ahora tomaban en cuenta la relación entre dos variables (con la excepción del coeficiente gamma). La correlación parcial provee medidas del grado de relación entre una variable dependiente Y y cualquiera de un conjunto de variables independientes, controladas por una o más de esas variables independientes. Es decir, describe la relación entre dos variables, controlando los efectos de una o más variables adicionales.

Es similar a lo que se hace en tabulaciones cruzadas, cuando se introducen variables de control. Sin embargo, ya habíamos visto que para controlar varias variables con varios valores necesitábamos una muestra demasiado grande, además de que la inspección del efecto era de tipo literal.

Con correlaciones parciales, el control no solamente es estadístico, sino que además la cantidad de casos no necesita ser muy grande.

$r_{ij \cdot k}$ indica entonces a i y j variable independiente y dependiente (el orden es

inmaterial, ya que la correlación entre ij y ji serán idénticas). La variable de control es indicada con k .³

Desde la perspectiva de la teoría de la regresión, la correlación parcial entre i y j , controlando por k , es la correlación entre los residuales de la regresión de i sobre k y de j sobre k , permitiéndonos establecer predicciones sobre las variables dependientes e independientes a partir del conocimiento del efecto que tiene la variable control sobre ellas.

El coeficiente de correlación parcial puede ser utilizado por el investigador para la comprensión y clarificación de las relaciones entre tres o más variables. Por ejemplo, puede ser utilizado para la determinación de espureidad, para la localización de variables intervinientes, y para la determinación de relaciones causales.

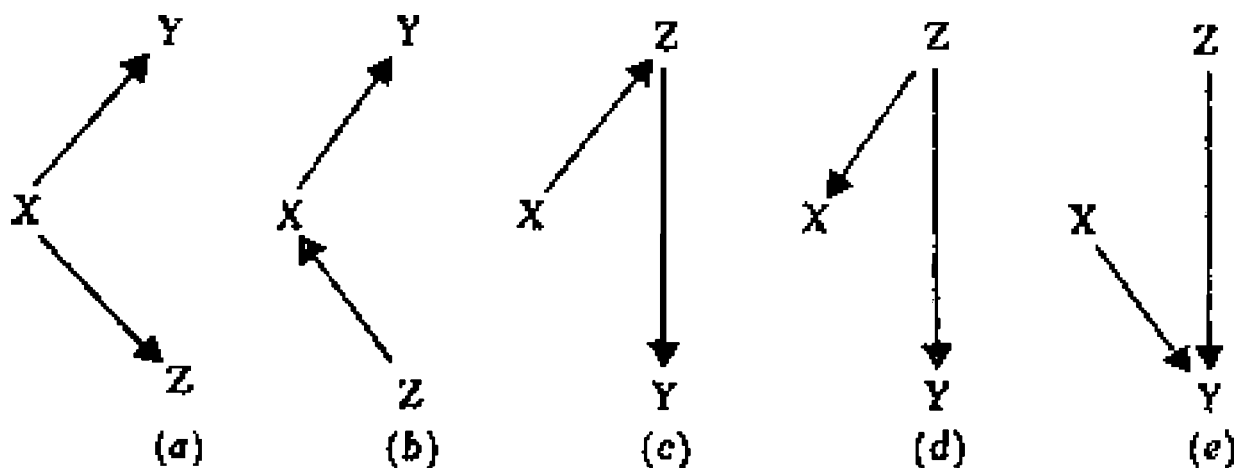
El coeficiente de correlación parcial *para la determinación de espureidad en las relaciones*: una relación espuria es aquella en la cual la correlación entre una variable X y una variable Y es el resultado de los efectos de otra variable (Z) que es el verdadero predictor de Y . La correlación es espuria cuando, controlando por Z (esto es, a Z constante), los valores de X no varían con los valores de Y . Éste es el caso en que los coeficientes de correlación parcial dan valores 0 o próximos a 0.

Supóngase una relación entre X y Y de .40. Computo un coeficiente de correlación parcial $r_{xy \cdot z}$ y el resultado es .20. Este coeficiente de correlación parcial ya me está indicando que la variable Z explica parcialmente la relación original entre X y Y . Computo un coeficiente de correlación parcial de segundo orden, en el cual controlo por dos variables, Z y W . El coeficiente de correlación parcial es ahora $r_{xy \cdot zw} = .06$, es decir que la relación desaparece, consecuentemente, la relación original era espuria.

Para la localización de variables intervinientes, así como para la determinación de relaciones causales, el problema es de naturaleza más conceptual, esto es, hay que combinar valores de coeficientes de correlación parcial con una serie de supuestos sobre las formas de las distribuciones y sobre la intervención de otras variables, además de las que se consideran en el modelo. Los supuestos no pueden ser verificados empíricamente por el análisis estadístico, sino que van a depender del razonamiento teórico.

En cualquiera de los siguientes casos, salvo en (e) la correlación parcial $r_{yz \cdot x}$ debe ser próxima a cero (Y es la variable dependiente, es decir, la que va a ocurrir al final en la secuencia temporal); (d) es un caso típico de correlación espuria. La relación entre X y Y se explica en función de las relaciones de X con Z y de Y con Z . (Véase la figura que viene a continuación.)

En el modelo (c) Z actúa como variable interviniente en la relación entre X y Y . La correlación parcial también dará 0. Pero hay que tener mucho cuidado en no interpretar los modelos (c) y (d) de la misma manera, y la correlación parcial tiene sentido solamente para probar que no hay relación entre X y Y , sino cuando interviene Z . El modelo (b) es similar, aunque ahora X es interviniente.



En el modelo (a) la relación X con Y y la X con Z son relaciones directas, mientras que no se postula relación entre Y y Z .

En los modelos (a) y (b), la correlación parcial entre X y Z , controlado por Y , debe ser 0.

Similarmente, en los modelos (c) y (d), la correlación parcial entre X y Y , controlado por Z , debe ser 0.

Cuando el modelo es (e), la correlación parcial $r_{xy \cdot z}$ dará valores más altos que la correlación entre X y Y . La correlación entre X y Z será 0.

El investigador debe informar a los programadores sobre la lista deseada de correlaciones parciales, en la que se especifiquen las combinaciones de variables (todas las combinaciones posibles o solamente algunas de las combinaciones). Por ejemplo: si se presentan variables: ingreso, educación, actitud frente al cambio y religiosidad, o se especifican las combinaciones deseadas o se deja que se correlacionen todas con todas en n combinaciones.

La palabra *with* especifica en el programa la combinación entre variables cuando la lista incluye solamente algunas combinaciones. Cuando el programa no incluye *with* se calculan todas las combinaciones posibles.

ANÁLISIS DE REGRESIONES MÚLTIPLES (SUBPROGRAM REGRESSION)

Este subprograma es considerablemente más complejo que los anteriores, y puede ser utilizado para una variedad bastante grande de análisis de variables múltiples: regresiones polinomiales, regresiones mudas (*dummy*), análisis de la varianza y análisis de la covarianza, predicciones, etcétera.

Por lo general, la regresión múltiple requiere variables medidas a nivel intervalar o racional y que las relaciones sean lineares y aditivas. Sin embargo, hay casos especiales en los cuales regresores mudos, medidos a nivel nominal, pueden ser incorporados a la regresión, relaciones no lineares y no aditivas pueden ser manipuladas, etcétera.

Existen algunas diferencias entre análisis de correlaciones múltiples y análisis de regresiones múltiples, que conviene destacar.

Los análisis de correlaciones múltiples se utilizan para: a) La evaluación de la medida

en que cada variable predictora o subconjunto de variables contribuye a la explicación de los puntajes de un criterio sobre una muestra; o *b*) Para predecir los puntajes de un criterio en una muestra diferente en la cual existe información del mismo grupo de variables predictoras. Aquí no estamos interesados tanto en la relación entre la variable dependiente y cada una de las variables independientes tomadas separadamente, sino en el poder explicativo del conjunto de variables independientes en su totalidad. El coeficiente de correlación múltiple es expresado entonces como $r_{1 \cdot 2345}$ si son 4 las variables predictoras, o $r_{1 \cdot 23}$ si son dos, etcétera.

Los modelos para el análisis de regresiones múltiples, a la vez que son más complejos en términos de cantidad de operaciones o de derivaciones que a través de ellos se puedan realizar, son bastante más simples en términos de los supuestos y condiciones para su utilización.

Por ejemplo, los modelos correlacionales requieren que las variables y los parámetros observables tengan una distribución normal conjunta; los modelos de regresión múltiple requieren solamente que la distribución de las desviaciones de la función de regresión sea normal; no se supone que las variables predictoras provengan de una distribución normal multivariata, o a veces requiere que los datos estén contenidos en códigos binarios.

Nosotros vamos a dar algunos ejemplos de prueba de hipótesis a través de análisis de regresiones múltiples. El análisis de las regresiones múltiples puede ser utilizado ya sea para la descripción de las relaciones entre variables o como instrumento para la inferencia estadística.

Como instrumento descriptivo la regresión múltiple es útil: *a*) Para encontrar la mejor ecuación lineal de predicción y para evaluar su eficiencia predictiva; *b*) Para evaluar la contribución de una variable o un conjunto de variables; *c*) Para encontrar relaciones estructurales y proveer explicaciones para relaciones complejas de variables múltiples.

Habíamos visto que el coeficiente de regresión simple se expresaba en la fórmula:

$$Y = \alpha + \beta X$$

Un coeficiente de correlación parcial, dijimos, era una medida de la cantidad de variación explicada por una variable independiente, después que las otras variables han explicado todo lo que podían. En el coeficiente de correlación múltiple estamos interesados en el poder explicativo de un conjunto de variables independientes sobre la variación dependiente ($r_{1 \cdot 2345}$).

Para ambos casos, la ecuación de la regresión toma ahora la siguiente forma:

$$Y = \alpha + \beta_1 X_1 + \beta_2 X_2 + \dots + \beta_k X_k$$

Ésta es la ecuación más simple, y parte de los mismos supuestos delineados para la ecuación de regresión simple. En la medida en que nos movemos en espacios multidimensionales, la representación geométrica es imposible.

Los coeficientes β se interpretan de manera distinta que en el caso de regresiones

simples, ya que aquí las inclinaciones son varias, y se obtienen cada una de ellas controlando por cada una de las variables independientes remanentes. Manteniendo X_2 en un valor fijo, β_1 representa la inclinación de la línea de regresión de Y sobre X_1 , para el caso en que solamente se estén controlando dos variables. Y así sucesivamente.

Ejemplos

a) Para encontrar la mejor ecuación lineal de predicción y para evaluar la eficiencia predictiva. Jae-On Kim y Frank Kohout⁴ presentan el problema de predecir la tolerancia política a partir de educación, ocupación e ingreso. A través de técnicas de regresión múltiple, el investigador podría estar interesado en determinar el grado de dependencia lineal de la tolerancia política sobre la base de la educación, la ocupación y el ingreso de una persona. Supóngase que la tabla de resultados es la siguiente:

Correlación múltiple	:	.5312
R^2	:	.2822
Error estándar	:	.8604
Variables independientes	B	β par
Educación	.1296	.3889
Ocupación	.0089	.1778
Ingreso	.0018	.0556
(Constante A)	2.9889	

La interpretación en este caso podría ser la siguiente: 1) La cantidad de variación en tolerancia política, explicada por la operación conjunta de educación, ocupación, e ingreso, es del 28.22% de la varianza total. 2) Si el investigador está interesado en predecir los puntajes que un sujeto va a obtener en tolerancia política a partir de las tres variables independientes, aplicará la ecuación de predicción señalada más arriba.

$$Y = 2.9889 + .1296 (X_1) + .0089 (X_2) + .0018 (X_3)$$

Si el sujeto tiene 10 años de educación formal (X_1), un puntaje de 60 en prestigio ocupacional (X_2) y un ingreso de 100 (\$ 10000) (X_3), entonces

$$Y = 2.9889 + .1296 (10) + .0089 (60) + .0018 (100) = 4.9989$$

El error estándar que figura en la tabla (.8604) predice que los puntajes precedidos en la escala de tolerancia política se van a desviar de los valores parámetros en .8604 unidades.

Los valores B en la tabla son coeficientes de regresión parcial, y pueden ser utilizados como medida de la influencia de cada variable independiente sobre la tolerancia política

cuando se controlan los efectos de las otras variables.

Obsérvese en el ejemplo que el coeficiente de correlación múltiple (R) es mayor en magnitud que cualquiera de los r , y esto es evidente desde el momento en que es imposible explicar menos variación agregando variables. El máximo valor relativo del coeficiente total ocurre cuando la correlación entre las variables independientes es igual a 0, de manera que, si queremos explicar la mayor cantidad posible de variación en la variable dependiente, deberemos buscar variables independientes que, si bien tienen correlaciones moderadas con la variable dependiente, son relativamente independientes unas de las otras.

Relacionado con la correlación entre variables independientes, está el problema de la *multicolinealidad*; esto es, cuando las variables independientes están estrechamente intercorrelacionadas, tanto las correlaciones parciales como la estimación de los β se hacen muy sensitivas a los errores de muestreo y de medición. Cuando la multicolinealidad es extrema (intercorrelaciones del rango de .8 a 1.0) el análisis de regresión no es recomendable.

b) La regresión múltiple puede ser utilizada también para evaluar la contribución de una variable independiente en particular, cuando la influencia de otras variables independientes es controlada. Aquí utilizamos coeficientes de regresiones parciales. Hay dos coeficientes designados, la contribución de cada variable a la variación de la variable dependiente: coeficiente de correlación semiparcial (*part-correlation*) y el coeficiente de correlación parcial. El primero se denota como $r_{y(1 \cdot 2)}$ y el segundo como $r_{(y1) \cdot 2}$.

El coeficiente semiparcial es la correlación simple entre el Y original y el residual de la variable independiente X^1 a la cual se le extraen los efectos de la variable independiente X^2 , es decir que el efecto de X^2 es sacado solamente de la variable X_1 , mediante una regresión linear simple de X_2 sobre X_1 , entonces ese residual de X_1 es correlacionado con la variable dependiente Y . En el caso de la tolerancia política uno podría estar interesado en determinar de qué manera el ingreso contribuye a la variación de la tolerancia política aparte de lo que es explicado por educación y ocupación. El cuadro siguiente permite calcular los valores del coeficiente semiparcial y el coeficiente parcial:

Regresión con dos variables independientes

A		B	C
Ed. (X_1) y Ocu. (X_2)		Ed. (X_1) e Ing. (X_3)	Ocu. (X_2) e Ing. (X_3)
Regress.	.5292	.5118	.4168
mult. (R)			
R^2	.2800	.2619	.1733

Regresión con tres variables independientes

Ed. (X_1) y Ocu (X_2), e Ing. (X_3)
Regresión múltiple (R) : .5312
 R^2 : .2822

Coefficiente semiparcial. Su cuadrado es igual a la diferencia entre un R^2 que incluye las tres variables independientes (.2822) y un R^2 que incluye solamente ocupación y educación (.2800). En nuestro caso entonces $R^2_{y(3.12)}$ es igual a .0022, indicando que ingreso solamente contribuye a un .22% de incremento en la variación de tolerancia política por educación y ocupación; en otras palabras, que el incremento es trivial, y que se puede ignorar ingreso. Para los casos del coeficiente semiparcial para educación y para ocupación los valores respectivos serían .1089 y .0203, es decir que educación explicaría aproximadamente un 11% de la variación y ocupación un 2%.

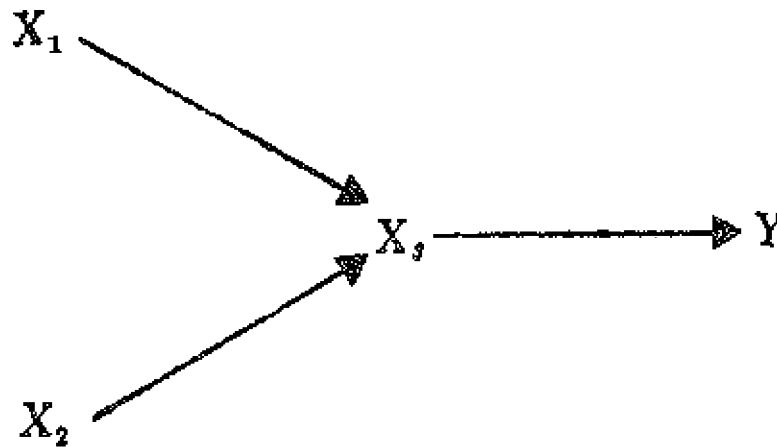
El coeficiente parcial. Es la correlación entre los dos residuales, el residual de Y y el residual de X_1 , para los cuales y en ambos se han extraído los efectos de X_2 . El cuadrado de una correlación parcial es el incremento proporcional en la variación explicada debido a X_1 , expresada como una proporción de la variación que no está explicada por X_2 . El coeficiente de correlación parcial indicaría el grado en que una variable da cuenta del remanente de variación del que no dan cuenta las otras variables independientes. En nuestro ejemplo de ingreso la correlación parcial es .0031, es decir que solamente da cuenta del .31% de la variable dependiente.

c) *El análisis de regresiones múltiples para la determinación de relaciones estructurales entre variables.* Se trata aquí de una conjunción de la técnica de regresión múltiple con la teoría causal. La teoría causal especificaría un ordenamiento de las variables que refleja una estructura de eslabones causa-efecto, la regresión múltiple determina la magnitud de las influencias directas e indirectas que cada variable tiene sobre las otras variables, de acuerdo al orden causal presumido. El método de *path analysis* es un método para descomponer e interpretar relaciones lineales entre conjuntos de variables, en los que se parte del supuesto de que el sistema causal es cerrado, consistente en causas y efectos encadenados. Las relaciones causales (*pathways*) se representan con flechas que conectan la causa al efecto.

Cuando se relacionan tres variables de las cuales una es dependiente (efecto), existen teóricamente seis maneras a partir de las cuales se puede establecer la relación (ver ejemplos en la sección de correlación parcial), con cuatro variables podemos producir 65 diferentes diagramas, etc. La tarea del investigador es seleccionar entre los diagramas posibles aquellos que sean más significativos desde el punto de vista de la teoría sustantiva. Cualquier diagrama, por ejemplo, el de la parte superior de la página siguiente, puede ser representado e interpretado en términos de ecuaciones estructurales: una variable a la que una o más flechas apuntan es interpretada como una función de solamente aquellas variables desde donde partan las flechas.

Uno de los supuestos principales del *path analysis* es que todas las relaciones son

lineares, que las variables son aditivas y que las relaciones son unidireccionales. Cuando se cumplen esas condiciones, la función linear toma la forma:



$$X_0 = C_{01}X_1$$

Donde:

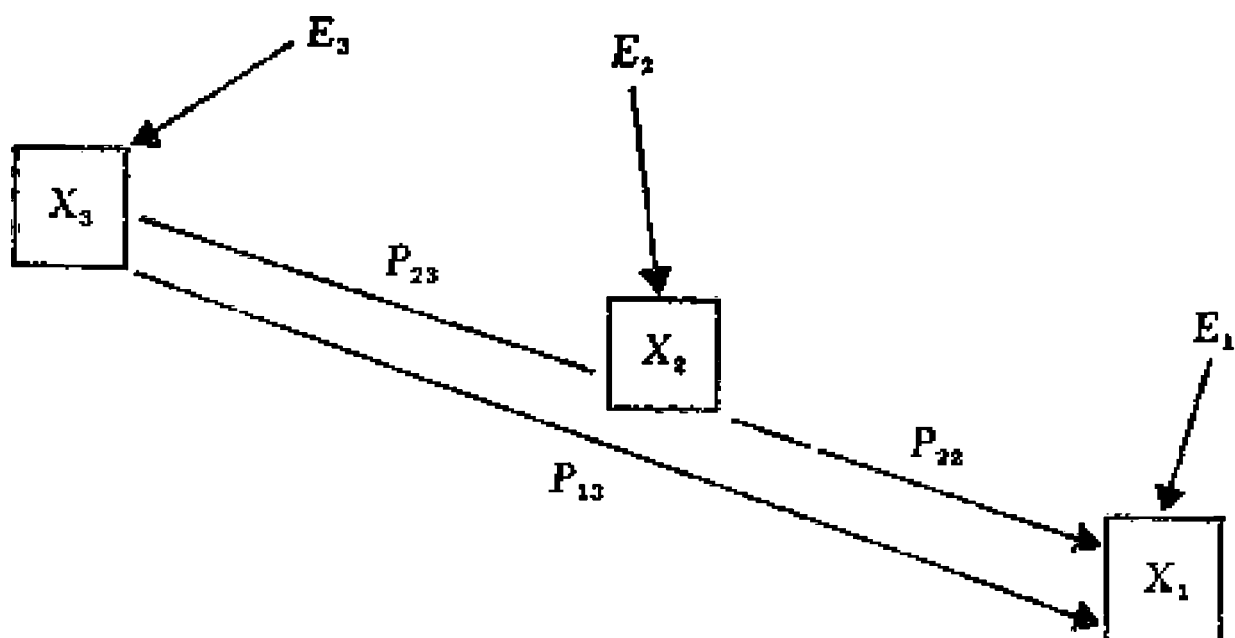
X_1 es la variable independiente, o causa,

X_0 es la variable dependiente, o efecto,

C_{01} es una constante que expresa la magnitud del cambio en X_0 para cada unidad de cambio en X_1 . Este coeficiente mide el efecto causal linear, o simplemente el coeficiente efecto.

El *path analysis* no es una técnica para demostrar causalidad. Es un procedimiento para el análisis de las implicaciones de un conjunto de relaciones causales que el investigador impone, a partir de algunos supuestos técnicos, en el sistema de relaciones.

Consideremos ahora un *path analysis* de tres variables, X_3 , X_2 , X_1 . Asumiendo que existe un orden en la relación entre las variables, digamos, $X_3 \geq X_2 \geq X_1$, y suponiendo que el sistema sea cerrado, podemos representar la relación de la siguiente manera:



o por un sistema de ecuaciones lineales tal como:

$$X_3 = E_3$$

$$X_2 = P_{23}X_3 + E_2$$

$$X_1 = P_{13}X_3 + P_{12}X_2 + E_1$$

$$\text{COV} (E_3, E_2) = \text{COV} (E_3, E_1) = \text{COV} (E_2, E_1) = 0$$

Cada E_i representa todos los efectos residuales en las causas de cada X_i y se denominan errores independientes o perturbaciones independientes, o variables latentes. Cada una de estas variables latentes se estima a partir de cada R^2 por medio de la fórmula $\sqrt{1-R^2}$, donde el coeficiente de correlación múltiple R es la parte de la ecuación de la regresión en la cual X_i es la variable dependiente y todas las variables que la causan son usadas como predictores.

Cada P_{ij} indica un *path* y puede ser estimado a partir de las regresiones de los X_i sobre los X_j . En el *path* que aparece más arriba P_{23} es estimado a partir de la regresión de X_2 sobre X_3 , donde $X_2 = B_{23}X_3$.

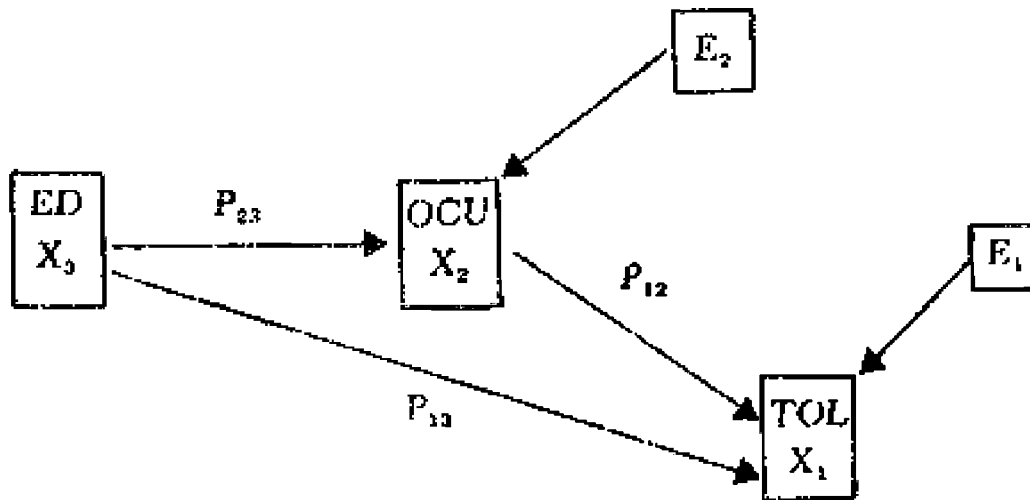
Y donde P_{13} y P_{12} pueden ser estimados de las regresiones de X_1 sobre X_2 y X_3 : $X_1 = B_{13}X_3 + B_{12}X_2$.

Por lo general, dadas n variables en orden $X_n \leq \dots \leq X_3, \leq X_2, \leq X_1$, la estimación de todos los *path* coeficientes requerirá $n - 1$ soluciones de regresión, en las que se toma cada una de las $n - 1$ variables de orden menor en el diagrama como independientes en sucesión y todas las variables de orden mayor como sus predictores.

Sigamos el mismo ejemplo de Kim y Kohout (*op. cit.*): tenemos 3 variables:

tolerancia (X_1), *status* ocupacional (X_2) y educación (X_3); si podemos sostener que el grado de tolerancia *probablemente* va a estar afectado por el nivel educacional y por el nivel del *status* ocupacional, y que el *status* ocupacional del individuo *probablemente* va a estar influido por su nivel educacional, entonces podemos postular un ordenamiento causal débil del tipo $X_3 \geq X_2 \geq X_1$. Estamos afirmando un juicio de probabilidad en el que no sabemos cómo una variable afecta a la otra. Para continuar con el *path analysis* necesitamos otro supuesto, que el sistema causal es cerrado, que es más difícil de justificar, pero supongamos que esté justificado.

Tenemos entonces un diagrama de la siguiente forma:



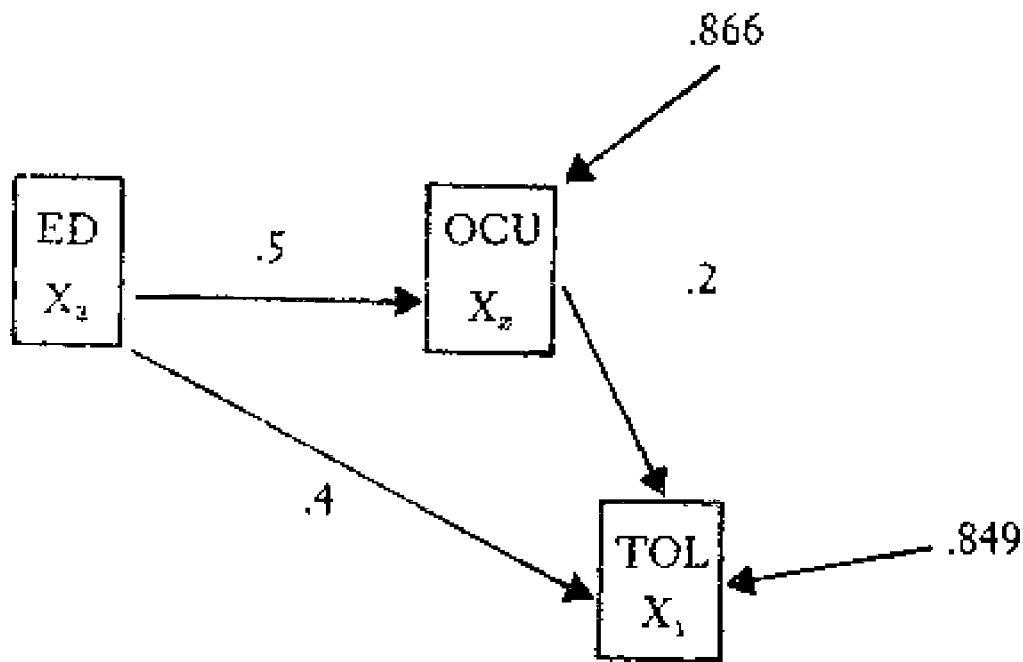
Calculamos los coeficientes de regresión simple para cada uno de los P_{ij} obteniendo los siguientes valores:

$$\begin{aligned}
 P_{23} &= .5 \\
 P_{12} &= .2 \\
 P_{13} &= .4
 \end{aligned}$$

Calculamos también los valores E_i que resultan ser:

$$\begin{aligned}
 E_2 &= .866 \\
 E_1 &= .8485
 \end{aligned}$$

El diagrama o *path analysis* tiene ahora la siguiente forma:



¿Cómo interpretamos este *path*? a) Primero examinamos cada subsistema, a través de las variables latentes. Y vemos que 75% de la variación en ocupación y 72% de la variación en tolerancia permanecen sin explicar por las relaciones causales explicitadas en el modelo.

b) Identificarnos los efectos de educación sobre ocupación; de ocupación sobre tolerancia, y de educación sobre tolerancia. El coeficiente C_{ij} mide los cambios que acompañan a X_i dada una unidad de cambio en X_j , estando controladas todas las causas extrañas. Los datos son los siguientes:

$$C_{23} = P_{23} = .5$$

$$C_{13} = (P_{23}) (P_{12}) + P_{13} = .5$$

$$C_{12} = P_{12} = .2$$

c) La covariación total entre pares de variables, representadas por la correlación simple, puede ser descompuesta de la siguiente manera:

	Ocu., Ed. (X_2, X_3)	Tol., Ed. (X_1, X_3)	Tol., Ocu. (X_1, X_2)
I) Covariación			
Original (r_{ij})	.5	.5	.4
II) b_1 : Causal-directa	.5	.4	.2
b_2 : Causal-indirecta	0	.1	0
Total causal (b_1) + (b_2) = c_{ij}	.5	.5	.2
III) No causal (A) - (B) = r_{ij} - c_{ij}	0	0	0

Para la relación entre ocupación y educación, el *path analysis* confirma los supuestos, todas las covariaciones entre los dos son tomadas como causales o genuinas.

La covariación entre educación y tolerancia es también tomada como causal, pero la covariación se descompone entre lo que es mediatizado por ocupación y entre lo que no lo es. Aquí parte de la relación entre educación y tolerancia está mediatizada por una variable interviniente.

La relación entre tolerancia y ocupación, esto es, la última columna, está descompuesta en componentes causales y componentes espurios.

Casos especiales en el path analysis

Hasta ahora consideramos modelos generales de *path analysis* en los cuales todas las relaciones bivariatas eran asumidas como teniendo una relación causal y el sistema como un todo era cerrado. Es posible introducir en el *path analysis* una cantidad de supuestos diferentes. Sin embargo, siempre hay que recordar que, cada vez que incorporamos supuestos ambiguos, producimos como resultado un modelo que da lugar a interpretaciones también ambiguas. Hasta ahora representamos las relaciones bivariatas como $X \rightarrow Y$, también podemos representar la relación entre las dos variables de las siguientes formas:

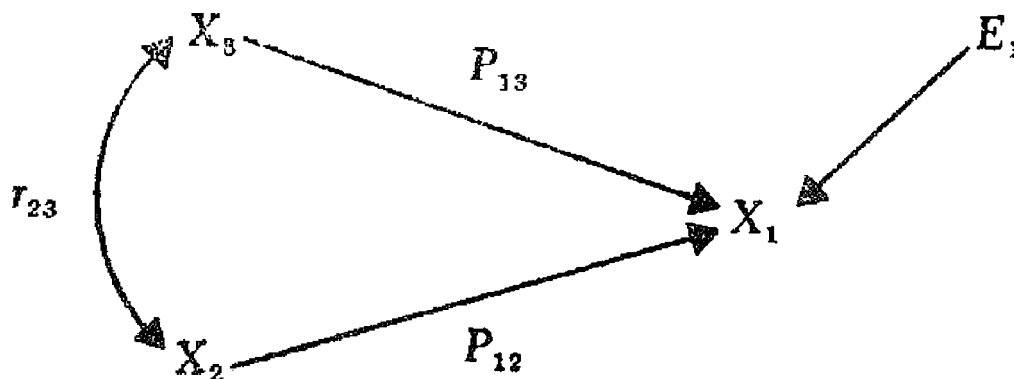
$X \curvearrowright Y$: Esto significa una correlación no analizada, por lo tanto la relación es ambigua, en el sentido en que la covariación puede ser causal o espuria, y la dirección de la relación puede ser de X a Y o de Y a X .

$X \quad Y$: La ausencia de flecha recta o curva significa que no existe covariación entre X y Y .

$X \curvearrowleft Y$: La curva simple significa que la relación entre X y Y es completamente espuria o no causal.

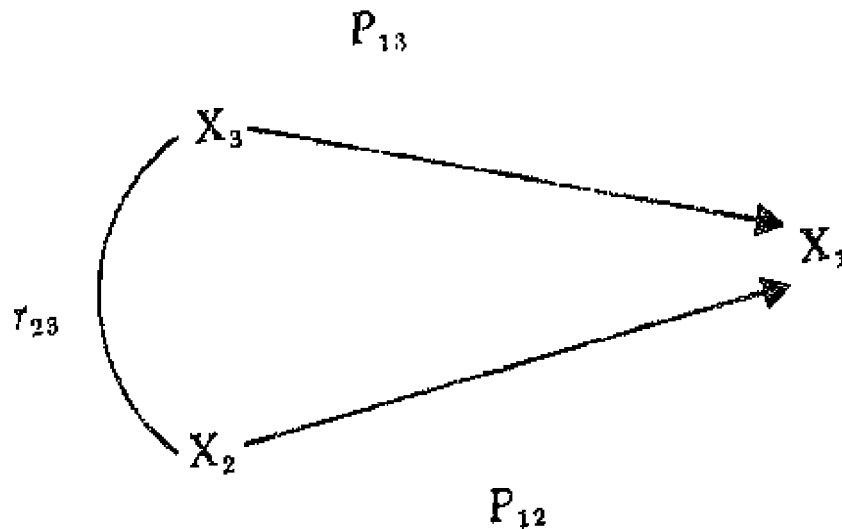
$X \rightrightarrows Y$: Representa una relación que es parcialmente causal y parcialmente espuria.

La relación representada hasta ahora en los diagramas anteriores era del tipo $X \rightarrow Y$, esto es, asumíamos un orden causal entre las variables. Existen situaciones en las cuales no conocemos la verdadera naturaleza de la relación causal entre algunas de nuestras variables aunque sí conocemos que existe correlación entre ellas. En este caso, el gráfico tendría la siguiente forma:



Es decir, postulamos relación causal entre X_3 y X_1 y entre X_2 y X_1 , pero las variables independientes no están conectadas entre sí por una conexión causal, sino simplemente por su correlación. La estimación de P_{13} y de P_{12} se obtiene a partir de las regresiones en las que X_1 es la variable dependiente y X_2 y X_3 , variables independientes. Las relaciones entre X_2 y X_3 se expresan por un coeficiente de correlación simple. Nótese que el cambio total en la variable dependiente no está definido en el modelo, lo que dificulta la predicción.

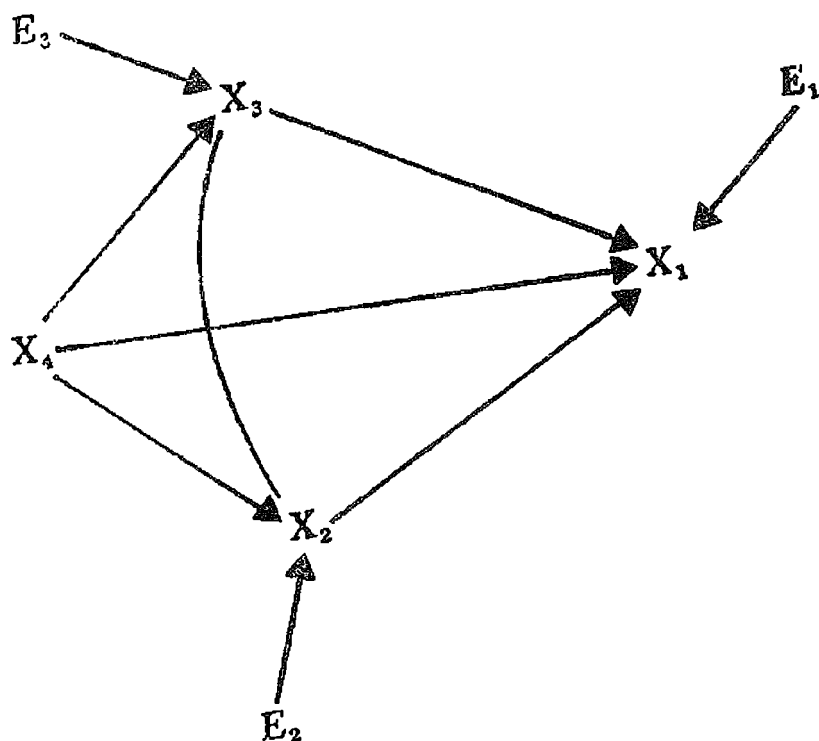
Cuando existen suficientes elementos en la teoría que efectivamente permiten asegurar que la covariación entre las variables exógenas no es de naturaleza causal, el modelo puede representarse de la siguiente forma:



Ahora sí todas las relaciones entre variables pueden interpretarse de manera causal, simplemente porque partimos del supuesto menos ambiguo que el del diagrama anterior. Aquí, en vez de plantear desconocimiento sobre la naturaleza de la covariación, planteamos que la covariación entre X_3 y X_2 es de naturaleza no causal, es decir que X_3 no causa variación en X_2 y viceversa. De esta manera es posible hacer predicciones en relación a los cambios que una unidad en X_3 o en X_2 producirán en X_1 .

Kim y Kohout presentan otro ejemplo con cuatro variables y un conjunto de supuestos fuertes, lo que da lugar a interpretaciones menos ambiguas: se trata de un esquema en el que se relaciona sexo (X_4) y accidentes de tráfico (X_1), controlando por cantidad de kilómetros conducidos al año (X_3) y frecuencia de conducción en horas de mucho tráfico (X_2).

Los supuestos causales son que el sexo puede afectar tanto la cantidad de kilometraje recorrido como las condiciones en las que se maneja, las cuales a su vez van a determinar las tasas individuales diferenciales de accidentes de tráfico. El investigador no tiene ningún supuesto teórico que le permita relacionar en forma causal el kilometraje recorrido con las condiciones de manejo. Asimismo, ni el total de kilometraje recorrido, ni las condiciones de manejo se considera que sean totalmente explicadas por sexo, sino que a su vez existen una serie de factores que pueden causar ambos. El modelo del *path* adquiere entonces la siguiente forma:



Donde las estimaciones de P_{34} y de P_{24} pueden realizarse a partir de coeficientes de correlación simple, y los P_{13} , P_{12} y P_{14} por las regresiones de X_1 , sobre X_2 , X_3 y X_4 .

El coeficiente de covariación residual entre X_2 y X_3 se obtiene a partir de: $r_{23} - (P_{34})(P_{24})$. Los coeficientes para los E_i se obtienen respectivamente de la siguiente forma:

$$E_1 = \sqrt{1 - R_{1.234}^2}$$

$$E_2 = \sqrt{1 - R_{2.34}^2}$$

$$E_3 = \sqrt{1 - R_{3.24}^2}$$

REGRESIONES CON VARIABLES MUDAS (DUMMY VARIABLES)

Éste es un caso especial de regresión, en el cual introducimos mediciones a nivel nominal en la ecuación de la regresión. Estas variables mudas se obtienen tratando cada categoría de la variable nominal como si fuera una variable por separado, asignando puntajes arbitrarios según la presencia o ausencia del atributo en cuestión. Por ejemplo,

si en afiliación política tenemos 3 partidos políticos: radical, demócrata cristiano y conservador, cada uno de esos partidos o categoría representa 1 de 3 variables dicotómicas, entonces los puntajes 1 a 0 pueden ser asignados a cada “variable”. Si un sujeto tiene afiliación radical, entonces su puntaje en radical será 1, su puntaje en demócrata cristiano será 0, y su puntaje en conservador será 0. Los valores 0 y 1 son tratados como variables intervalares e incluidas así en la ecuación de la regresión. Sin embargo, por un problema de álgebra ⁵ una de las variables mudas debe ser excluida de la ecuación de regresión. De hecho, la variable muda excluida actuará ahora como punto de referencia a partir del cual se interpretan los valores en cada una de las otras variables mudas. Cada categoría ahora es representada por una combinación de las i variables mudas. Supongamos en nuestro caso que la categoría de referencia es otro partido; tendríamos entonces la siguiente distribución de puntajes:

	X_1	X_2	X_3
Radical	1	0	0
Demócrata cristiano	0	1	0
Conservador	0	0	1
Otro	0	0	0

Si ha sido elegido otro como categoría de referencia, la ecuación de la regresión puede escribirse entonces como:

$$Y = A + B_1X_1 + B_2X_2 + B_3X_3$$

donde los casos de la categoría “otros” pueden predecirse mediante:

$$Y = A$$

los radicales por:

$$Y = A + B_1X_1$$

y en la medida en que el valor de radicales X_1 es 1, entonces:

$$Y = A + B_1$$

Los valores esperados para cada una de nuestras categorías serán entonces:

	<i>Y</i>
Radical	$A + B_1$
Demócrata cristiano	$A + B_2$
Conservador	$A + B_3$
Otro	A

Análisis de la varianza unidireccional con variables mudas

El análisis de la varianza unidireccional puede obtenerse a través de diferentes subprogramas en el SPSS: los subprogramas *anova*, *oneway* y *breakdown* (ver sección análisis de la varianza). En estos tres subprogramas las variables entran como variables nominales, no introduciéndose la creación de variables mudas.

Sin embargo, el investigador puede desear un análisis de la varianza unidireccional con el subprograma *regression*. Para ello debe crear un conjunto de variables mudas según el sistema explicado más arriba e instruir al programador para que incluya las instrucciones pertinentes para la creación de variables mudas en el SPSS.

El *out-put* del subprograma *regression* en su porción referente al análisis de la varianza unidireccional tiene la siguiente forma, en la cual introducimos cálculos ficticios para las tres variables usadas en nuestro ejemplo con la variable dependiente, actitud frente a la nacionalización del petróleo:

<i>R</i> múltiple	.5844	Análisis de var.	<i>D. F.</i>	<i>Suma cuadr.</i>	<i>F</i>
<i>R</i> ²	.3416	Regresión	3	56.8529	16.5993
Error estándar	1.0638	Residual	96	108.6371	

Variable en la ecuación

<i>Variable</i>	<i>B</i>	<i>Beta</i>	<i>Error estándar B</i>	<i>F</i>
D. cristiano	1.3156	.4435	.4135	10.121
Radical	-.3961	-.1497	.3795	1.089
Conservador	-.9411	-.1441	.6393	2.183
(Constante)	2.444			

El valor *F* de 16.5993 tiene una probabilidad mayor que .001, es decir que las diferencias son muy significativas para el conjunto de partidos. El *R*² es equivalente al coeficiente de correlación múltiple que se deriva del coeficiente de correlación eta (ver correlación simple), y su valor indica que el 34% de la actitud frente a la nacionalización

del petróleo depende o se explica por la afiliación política.

Los promedios para cada categoría pueden obtenerse a partir de la columna *B* de “variables en la ecuación” que es el *out-put* de la regresión:

— Actitud frente a la nacionalización del petróleo: $Y = 2.444$

— Radical: $Y = 2.444 + 1.3156 = 3.76$

— D. cristiano: $Y = 2.444 + (-.3961) = 2.05$

— Conservador: $Y = 2.444 + (-.9444) = 1.50$

Regresiones con variables mudas para dos o más variables categorizadas

La ecuación de predicción para dos variables nominales (representadas por dos conjuntos de variables mudas) es la siguiente:

$$Y = A + B_1X_1 + B_2X_2 + B_3X_3 + B_4E_1$$

Donde los X_1 representan una variable nominal con 4 categorías y E_1 una categoría de una variable nominal dicotómica.

El valor predictivo para cada celda de la matriz estará dado por el siguiente cuadro, siguiendo nuestros ejemplos anteriores:

	<i>Varón</i>	<i>Mujer</i>
Radical	$A + B_1 + B_4$	$A + B_1$
D. cristiano	$A + B_2 + B_4$	$A + B_2$
Conservador	$A + B_3 + B_4$	$A + B_3$
Otro	$A + B_4$	A

Lo que ocurre ahora es que las categorías mujer y otro actúan como categorías de referencia.

Análisis de la varianza multidireccional con variables mudas

La regresión múltiple con n variables mudas puede ser utilizada para computar análisis de la varianza. Cuando se desea computar análisis de la varianza con las variables nominales sin recurrir a variables mudas se recomienda el subprograma *anova* (ver más adelante en “análisis de la varianza”).

Cuando se utilizan variables mudas y queremos realizar análisis de la varianza con el subprograma *regression* es necesario agregar para el caso de dos factores (afiliación política y sexo) los efectos de interacción, es decir, necesitamos crear tres nuevas variables mudas en nuestro ejemplo: (X_1E_1) , (X_2E_1) , (X_3E_1) , donde la ecuación de la regresión múltiple tendrá ahora la siguiente forma:

$$Y = A + B_1X_1 + B_2X_2 + B_3X_3 + B_4E_1 + B_5(X_1E_1) + B_6(X_2E_1) + B_7(X_3E_1)$$

Esta regresión representa el *modelo saturado* donde todos los términos de interacción posible están incluidos.

Los valores predictivos para el modelo saturado se obtienen del siguiente cuadro:

	<i>Varón</i>	<i>Mujer</i>
Radical	$A + B_1 + B_4 + B_5$	$A + B_1$
D. cristiano	$A + B_2 + B_4 + B_6$	$A + B_2$
Conservador	$A + B_3 + B_4 + B_7$	$A + B_3$
Otro	$A + B_4$	A

Para dos variables nominales A y B , la estrategia del análisis de la varianza sigue lo que se llama modelo clásico de análisis de la varianza en el cual ni los factores A y B (afiliación política y sexo) son ortogonales; esto es, si las frecuencias en las celdas son proporcionales a las frecuencias marginales de afiliación política y sexo, la suma de $2a$ y de $2b$ será simplemente la suma de los cuadrados debidos a cada factor y será igual a la suma de los cuadrados debidos a efectos aditivos. Si A y B no son ortogonales, los efectos de A se confundirán con los efectos de B , y la suma de $2(a)$ y de $2(b)$ no será igual a la suma de los efectos aditivos. El siguiente cuadro ilustra el modelo clásico:

<i>Fuente de variación</i>	<i>Suma de cuadrados</i>	<i>Df</i>	<i>F</i>
1) Suma de cuadrados debido a <i>A</i> y <i>B</i> modelo saturado	$SS_y(R_{A, B, AB}^2)$	$K = k_1 + k_2 + k_1 k_2$	$\frac{(1)/K}{(4)/(N-K-1)}$
2) Suma de cuadrados debidos a <i>A</i> y <i>B</i> modelo aditivo	$SS_y(R_{A, B}^2)$	$k_1 + k_2$	$\frac{(2)/(k_1 + k_2)}{(4)/(N-K-1)}$
a) Suma de los cuadrados debidos a <i>B</i> ajustados por <i>A</i>	$SS_y(R_{A, B}^2 - R_A^2)$	k_1	$\frac{(2a)/k_1}{(4)/(N-K-1)}$
b) Suma de los cuadrados debidos a <i>A</i> ajustados por <i>B</i>	$SS_y(R_{A, B}^2 - R_B^2)$	k_2	$\frac{(2b)/k_2}{(4)/(N-K-1)}$
3) Suma de los cuadrados debida a la interacción	$SS_y(R_{A, B, AB}^2 - R_{A, B}^2)$	$k_1 k_2$	$\frac{(3)/k_1 k_2}{(4)/(N-K-1)}$
4) Suma de los cuadrados residuales	$SS_y(1 - R_{A, B, AB}^2)$	$N - K - 1$	

Los significados de este cuadro serán analizados en la sección siguiente que corresponde a análisis de la varianza.

ANÁLISIS DE LA VARIANZA Y DE LA COVARIANZA (SUBPROGRAMAS ANOVA Y ONEWAY)

El análisis de la varianza es una técnica estadística utilizada para la determinación de asociación entre dos o más variables. El análisis de la varianza simple (o unidireccional) se refiere a situaciones en las cuales el investigador está interesado en determinar los efectos de una variable o factor (medido a nivel nominal) sobre una variable dependiente (o variable criterio) que debe estar medida a nivel intervalar. Si el investigador está interesado en el *efecto simultáneo de varios factores*, entonces el análisis de la varianza es bivariato o n-maneras.

En el análisis de la covarianza el investigador está interesado en los efectos tanto de variables no métricas como de variables métricas.

Análisis de la varianza simple

Habíamos visto en la sección correspondiente a confiabilidad de la diferencia entre

estadísticos algunas pruebas para la determinación de diferencias significativas en pares de medias muestrales (prueba t de Student). El análisis de la varianza simple, sobre todo en lo referido a la prueba F de Fisher, se emplea de manera similar, aunque ahora sea para decidir estadísticamente si la serie de datos entre n pares de medias son lo suficientemente diferentes entre sí para permitirnos el rechazo de la hipótesis de que esas medias surgieron, por efectos del azar, de una población única.

La *varianza total* en muestras combinadas tiene dos componentes: un componente representado por la varianza *interserial* (es decir, por la suma de los desvíos al cuadrado de las medias de las submuestras con respecto a la media total), y un componente representado por la varianza *intraseriale* (es decir, por la suma de los desvíos al cuadrado dentro de cada una de las series de datos).

Es decir que la varianza total se descompone en dos componentes: una intervarianza y una intravarianza: la prueba F es la razón entre la intervarianza y la intravarianza:

$$F = \frac{\text{intervarianza}}{\text{intravarianza}}$$

La intervarianza se estima sobre la base de k medias, las que pueden ser consideradas como k datos independientes. Puesto que se quiere una estimación de la varianza de la población, la suma de los desvíos al cuadrado se divide por los grados de libertad. Es decir, en el caso de la intervarianza sería $k - 1$. La fórmula para el cálculo de la intervarianza será entonces:

$$\text{Intervarianza} = \frac{\sum n_i d_i^2}{k - 1}$$

Donde:

n = Cantidad de casos en cada una de las series.

d = Desvío de la media serial con respecto a la media total ($M_g - M_t$).

k = Cantidad de series.

La intravarianza se estima sobre la base de las medias muestrales o seriales, de donde:

$$\text{Intravarianza} = \frac{\sum x_i^2}{k (n - 1)}$$

Donde:

x = Desviación del puntaje de su media muestral.

La confrontación de los valores F se hace con tablas especiales, que señalan el nivel de significación. Las F mayores o iguales a los distintos niveles indican el grado de confianza con el cual se puede rechazar la hipótesis nula. La prueba F indica solamente

si existe o no una diferencia. Para encontrar dónde se encuentra ésta se hace necesaria una investigación y *tests* sucesivos, por ejemplo, pruebas *t*.

Un ejemplo de análisis de la varianza simple. Se trata de saber si el nivel de ingresos difiere por tipo de rama ocupacional. La variable dependiente o criterio será “ingreso” en escudos (los datos corresponden a Chile en 1968), por lo tanto, la variable está medida a nivel racional. La variable independiente o factor es sector, distinguiendo tres categorías: industria, construcción y servicios.

El siguiente cuadro muestra la distribución:

CUADRO 1. Distribución de ingreso por tipo de ocupación

	<i>Industria</i>	<i>Construcción</i>	<i>Servicios</i>
0- 100	7	1	4
100- 200	12	4	6
200- 300	34	12	19
300- 400	62	30	31
400- 500	51	24	24
500- 600	20	12	12
600- 700	22	15	8
700- 800	21	8	9
800- 900	17	8	7
900- 1000	9	5	5
1000- 1400	5	5	8
1400- 1800	8	4	6

La hipótesis nula a probar es que no existen diferencias entre los distintos promedios de ingresos según las tres ramas ocupacionales.

Los resultados obtenidos son los siguientes:

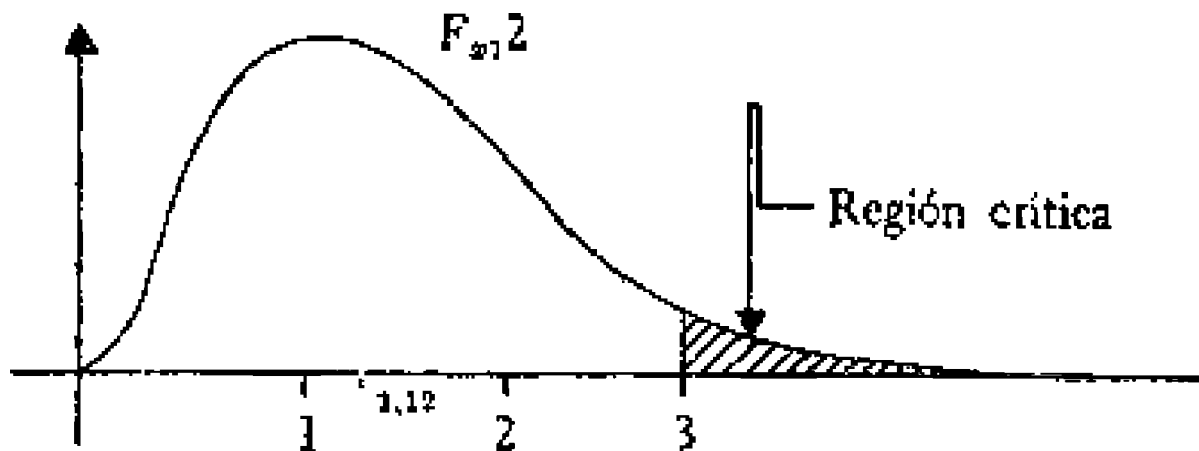
— Suma de las desviaciones al cuadrado de los datos individuales respecto a la media total	53 471 950
— Suma de las desviaciones al cuadrado de los datos individuales respecto a la media serial	178 652
— Suma de las desviaciones al cuadrado de las medias seriales respecto a la media total	53 293 342
Estimación de la intervarianza:	100 175
Estimación de la intravarianza:	89 326

$$F = 1.12$$

La distribución para *F* con 2 grados de libertad en el numerador y 532 en el denominador nos dice que hay una probabilidad de 5% de obtener un valor de *F* más alto

que 3. (3)

FIGURA 2



Como el valor obtenido para F cae fuera de la región crítica, no podemos rechazar la hipótesis nula de que las medias de las poblaciones son iguales. Por lo tanto, se concluye que los distintos sectores de la economía considerados no difieren en los niveles de ingreso de los trabajadores.

De haber sido F lo suficientemente grande como para caer en la región crítica, se hubiera rechazado la hipótesis nula, correspondiendo un análisis ahora de las diferencias entre pares de medias seriales.

El subprograma *breakdown* computa análisis de la varianza simple o unidireccional, computando medias, desviaciones estándar, intervarianza, intravarianza y valores F .

Ya vimos más arriba que el subprograma *regression* computaba también análisis de la varianza, aunque son variables mudas.

El subprograma *oneway* computa también análisis de la varianza simple; esto es, con una sola variable independiente o factor. La ventaja de este subprograma sobre los anteriores es que acepta hasta un máximo de 20 variables dependientes; claro que hay que especificar una sola variable independiente, y los *out-puts* son de dos en dos.

Análisis de la varianza n-dimensional

En el análisis de la varianza simple las series de datos se diferenciaban en base a un solo factor. En un problema de clasificación bidireccional existen dos bases distintas para la clasificación; en un análisis tridireccional habrá tres bases distintas para la clasificación, etcétera.

Un problema de análisis bidireccional típico puede ser la comparación de distintos métodos para la enseñanza de lectura y escritura en la escuela primaria, combinados con distintos tipos de maestros. A este respecto, la comparación de cuatro métodos (global, palabra generadora, silábico y "novo") con 3 clases distintas de maestros (normales fiscales, normales privados, especiales) genera 12 combinaciones posibles de método-maestro. Si se quisiera incluir una variable o factor adicional, digamos sexo, las

combinaciones posibles aumentarían a 24.

Las fuentes de la varianza en una clasificación bidireccional son ahora: *a)* Una varianza relativa al tipo de método de enseñanza; *b)* Una varianza relativa al tipo de maestro; *c)* Una varianza de la interacción entre determinado tipo de maestro y determinado tipo de método de enseñanza, y *d)* Una varianza residual o intravarianza, que constituye la estimación básica o residual de la varianza una vez que las tres fuentes de variación han sido eliminadas. Esta varianza residual puede ser considerada entonces como una varianza del error, puesto que representa las influencias de los factores desconocidos o no controlados.

Los grados de libertad en clasificaciones bidireccionales dependen de la fuente u origen de la varianza y son:

Entre líneas	$r - 1$
Entre columnas	$k - 1$
Interacción	$(r - 1) (k - 1)$
Dentro de las submuestras	$rk (n - 1)$
Total	$N - 1$

El *out-put* del subprograma *anova* provee para un cuadro como el ilustrado los siguientes datos de análisis de la varianza:

Efectos principales	Suma de cuadrados	DF	F	Nivel de sign.
Métodos	-----	---	-	-----
Maestros	-----	---	-	-----
Método-maestro	-----	---	-	-----
Residual	-----	---	-	-----
Total	-----	---	-	-----

Para tablas más complejas (digamos, tres variables factor), el programa proveerá los cálculos correspondientes a los efectos principales de cada uno de los efectos, 3 efectos de interacción bivariata por ejemplo, método-maestro, método-sexo, maestro-sexo; el efecto de interacción de los tres factores (método-maestro-sexo), la varianza residual y la varianza total.

Acorde con los valores, el investigador acepta o rechaza sus hipótesis nulas (relativas a diferencia entre métodos de enseñanza, tipo de maestros, sexo de los mismos; además de las interacciones método-maestro; método-maestro-sexo, etc.). Hay que recordar que, para la selección de las combinaciones o de métodos o de maestros especiales, una vez que se encuentran valores *F* significativos, es necesario aplicar pruebas *t*, con el fin de seleccionar aquellos que son más eficaces para el aprendizaje de los niños.

El cuadro 2 resume las distintas fuentes de variación para problemas *n*-direccionales.

En el tipo de ejemplo que describimos (método-maestros), la matriz de datos a partir

de la cual se realiza el análisis de la varianza se representa en el cuadro 2.

El subprograma *anova* permite cálculos de análisis de la varianza para un máximo de 5 factores o variables independientes en cada diseño. Como opciones permite análisis de la covarianza (hasta 5 covariaciones) y un cuadro de análisis de clasificaciones múltiples.

El subprograma *anova* es aplicable tanto para diseños ortogonales (igual número de frecuencias en cada una de las celdas) como para diseños no ortogonales (distintos números de frecuencias en las celdas). Incluso el programa puede considerar algunas celdas vacías.

El subprograma *anova* también puede producir un cuadro de *análisis de clasificación múltiple*, por medio del cual los resultados del análisis de la varianza son expuestos de manera más específica. Este método es particularmente útil cuando los efectos de interacción no son significativos y cuando los factores son variables nominales o atributos que no han sido manipulados experimentalmente, y por lo tanto pueden estar intercorrelacionados. Para dos o más factores interrelacionados puede ser importante conocer el efecto neto de cada variable, cuando se controlan los otros factores. Seguimos el ejemplo de Kim y Kohout (*op. cit.*) para ilustrar la interpretación de un cuadro completo de análisis de clasificación múltiple:

La variable dependiente es el salario semanal de los empleados de una industria. Los factores son sexo y raza. En la medida en que se sospecha cierto grado de discriminación social, el investigador está interesado en los efectos de los factores raza y sexo. Se sabe que dos variables adicionales, nivel de educación y duración en el empleo, determinan el nivel de los salarios, así que ambas son introducidas como variables. Los resultados obtenidos son los siguientes:

Clasificación de análisis múltiple: salario por sexo, raza, con educación y duración en el empleo.

Media total: 100 dólares

CUADRO 2

Líneas (maestros)	Sujetos	Columnas (medidas)				Suma de las líneas	Medias de las líneas
		1	2	3	4		
A	1	X_{a1}	X_{a2}	X_{a3}	X_{a4}		
	2	.	.	.	.		
	3	.	.	.	.		
	4	.	.	.	.		
	.	.	.	.	.		
	.	.	.	.	.		
	Σ M	ΣX_{a1} M_{a1}	ΣX_{a2} M_{a2}	ΣX_{a3} M_{a3}	ΣX_{a4} M_{a4}	ΣX_a	ΣM_a
B	1	X_{b1}	X_{b2}	X_{b3}	X_{b4}		
	2	.	.	.	.		
	3	.	.	.	.		
	4	.	.	.	.		
	.	.	.	.	.		
	.	.	.	.	.		
	Σ M	ΣX_{b1} M_{b1}	ΣX_{b2} M_{b2}	ΣX_{b3} M_{b3}	ΣX_{b4} M_{b4}	ΣX_b	ΣM_b
C	1	X_{c1}	X_{c2}	X_{c3}	X_{c4}		
	2	.	.	.	.		
	3	.	.	.	.		
	4	.	.	.	.		
	.	.	.	.	.		
	.	.	.	.	.		
	Σ M	ΣX_{c1} M_{c1}	ΣX_{c2} M_{c2}	ΣX_{c3} M_{c3}	ΣX_{c4} M_{c4}	ΣX_c	ΣM_c
Suma de las colum- nas (ΣX_k)		ΣX_1	ΣX_2	ΣX_3	ΣX_4	ΣX_k	M_T
Medias de las co- lumnas (M_k)		M_1	M_2	M_3	M_4		Media To- tal

La primera columna del cuadro de la página siguiente expresa las medias en cada categoría como desviaciones de la media total.

Las medias en la segunda columna expresan las medias de cada categoría (también como desviaciones de la media total), pero ahora ajustadas. Las disminuciones en los

valores indican que, en el contexto del empleo, sexo y raza están relacionados (cada factor disminuye cuando lo ajustamos por el otro factor). Los valores indican que los empleados varones tienden a ser blancos, mientras que las empleadas mujeres tienden a ser no blancas.

<i>Variables</i>	<i>Desviaciones de la media total</i>		
	<i>No ajustada</i>	<i>Ajustada por independientes</i>	<i>Ajustada por independientes y covariables</i>
Raza			
1—Blancos	+10	6	4
2—No blancos	-40	-24	-16
(eta y beta)	(.632)	(.384)	(.258)
Sexo			
1—Varones	12	8	6
2—Mujeres	-18	-12	-9
(eta y beta)	(.465)	(.310)	(.232)
<i>R</i> Múltiple	...	.648	.866
<i>R</i> ²	...	.620	.750

La tercera columna indica que, cuando se introduce educación y duración en el empleo, la influencia del sexo y de la raza disminuye, aunque aún persiste la discriminación.

Los beta y eta que figuran en el cuadro ayudan a su mejor interpretación. Los valores de la primera columna son valores eta, mientras que los de la segunda y la tercera son beta parciales. Comparándolos vemos que, tanto para el caso de sexo como en raza (en este último con mayor intensidad), la relación decrece a medida que se introducen más variables de control.

La correlación múltiple indica la relación entre salario y, como variable dependiente, los efectos aditivos de sexo y raza en la segunda columna; y de sexo, raza, educación y antigüedad en la última columna.

ANÁLISIS FACTORIAL (SUBPROGRAMA FACTOR)

El análisis factorial es una técnica matemática cuyo objetivo más amplio es el descubrimiento de las dimensiones de variabilidad común existentes en un campo de fenómenos. Cada una de estas dimensiones de variabilidad común recibe el nombre de *factor*. El razonamiento subyacente es el siguiente: si tenemos un conjunto de fenómenos, y si cada fenómeno varía independientemente de los demás, entonces habrá tantas dimensiones de variación como fenómenos. Por el contrario, si los fenómenos no varían independientemente, sino que hay ciertas dependencias entre ellos, entonces

encontraremos que las dimensiones de variación serán menores que los fenómenos. El análisis factorial, a través de una serie de procedimientos, nos permite detectar la existencia de ciertos patrones subyacentes en los datos de manera que éstos puedan ser reagrupados en un conjunto menor de factores o componentes.

Hay cuatro pasos fundamentales en el análisis factorial: *A)* Preparación; *B)* Factorización; *C)* Rotación; *D)* Interpretación.

Dentro de cada uno de los pasos existen procedimientos u opciones que se irán detallando en la medida en que desarrollemos cada uno de esos pasos.

A) Preparación

Consiste tanto en el planteo del problema a tratar cuanto en la formulación de hipótesis y recolección de datos. El tipo de variables que el investigador utilice tendrá importancia fundamental tanto en lo referente a los factores como a la interpretación. La mayoría de las técnicas analíticas requieren variables intervalares al menos, aunque es posible utilizar algunas de las medidas de asociación que discutimos en otras secciones de este trabajo. Lo importante es que el resultado de este primer paso de preparación es una matriz de correlaciones, que adquiere ya sea la forma de un triángulo o de un cuadrado:

		1	2	3	4	5	6	7	8		n
V a r i a b l e s	1	r_{11}	r_{12}	r_{13}	r_{14}	r_{15}	r_{16}	r_{17}	r_{18}		r_{1n}
	2		r_{22}	r_{23}	r_{24}	r_{25}	r_{26}	r_{27}	r_{28}		r_{2n}
	3			r_{33}	r_{34}	r_{35}	r_{36}	r_{37}	r_{38}		r_{3n}
	4				r_{44}	r_{45}	r_{46}	r_{47}	r_{48}		r_{4n}
	5					r_{55}	r_{56}	r_{57}	r_{58}		r_{5n}
	6						r_{66}	r_{67}	r_{68}		r_{6n}
	7							r_{77}	r_{78}		r_{7n}
	8								r_{88}		r_{8n}
	.										
	.										
	.										r_{nn}

El otro lado del triángulo para completar el cuadrado representa los mismos números o correlaciones, ya que la correlación r_{12} es idéntica a r_{21} , etc. La diagonal designa la correlación de la variable consigo misma.

El investigador tiene una opción en términos de preparación de la matriz de

correlaciones: se trata de lo que se da en denominar *Q-factor analysis* o *r-factor analysis*.

Si el análisis factorial se aplica a la matriz de correlaciones de *unidades*, donde por unidad se entiende el objeto, persona, etc., que detenta una característica o conjunto de características (es decir, donde correlacionamos pares de unidades), entonces se trata de *Q-factor analysis*. Si las correlaciones se hacen entre *variables* entre cada par de características o atributos, la técnica se denomina *R-factor analysis*.

En términos de preparación de matriz, los usuarios del SPSS en realidad no tienen esta opción, ya que el subprograma *factor* únicamente opera con el análisis de tipo *R*. Lo importante a destacar en esta sección es que el subprograma acepta como *input* datos brutos, matrices de correlaciones o matriz factorial.

B) Factorización

La factorización trata de poner de manifiesto por métodos matemáticos cuántos factores comunes es preciso admitir para explicar los datos originales o la matriz de intercorrelaciones.

Por este procedimiento surgen “nuevas variables” o factores que pueden ser definidos como transformaciones matemáticas exactas de los datos originales (análisis de componentes principales), o a través de supuestos inferenciales acerca de la estructura de las variables y de su fuente de variación (análisis factorial clásico o de factores inferidos). Ya sea que los factores sean definidos o inferidos, los factores iniciales son extraídos de tal manera que sean independientes los unos de los otros, esto es, factores que sean ortogonales. En este segundo estadio importa más la reducción de la matriz o de dimensiones que la localización de dimensiones significativas.

El análisis de los componentes principales no requiere ningún supuesto acerca de la estructura subyacente al conjunto de variables. Simplemente, trata de encontrar la mejor combinación lineal de variables tal que dará cuenta de una mayor proporción de la varianza que cualquier otra combinación lineal posible. El primer componente principal es entonces el mejor conjunto de relaciones lineales entre los datos; el segundo componente es la segunda combinación lineal tal que no está correlacionada con el primer componente (es decir, es ortogonal al primer componente); el segundo factor da cuenta de la varianza residual no explicada por el primer factor; el resto de los componentes es definido en forma similar, siendo los componentes tantos hasta cuando se haya explicado totalmente la varianza.

El modelo del componente principal puede ser expresado como:

$$z_j = a_{j1}F_1 + a_{j2}F_2 + a_{j3}F_3 + \dots + a_{jn}F_n$$

Donde:

a_{ji} = Coeficiente de regresión múltiple estandarizado de la variable j sobre el factor i ;

F_i = Factores definidos.

El análisis de factores inferidos está basado en el supuesto de que las correlaciones empíricas son el resultado de alguna regularidad subyacente a los datos. Se supone que cada variable está influida por varios determinantes, algunos de los cuales son compartidos por otras variables (determinante común) y por otros determinantes que no son compartidos por ninguna de las otras variables en el modelo (determinante único). Se supone entonces que las correlaciones entre variables son el resultado de variables compartiendo determinantes comunes. Se espera, por consiguiente, que el número de determinantes sea menor que el número de variables.

El modelo se expresa ahora de la siguiente forma:

$$z_j = a_{j1}F_1 + a_{j2}F_2 + \dots + a_{jm}F_m + d_jU_j$$

Donde:

U_j = Factor único para la variable j ;

d_j = Coeficiente de regresión estandarizado de la variable j sobre el factor único j .

Se asume en el modelo que $r_{(z_j, v_j)} = 0$, y que $r_{(v_j, v_k)} = 0$, es decir que el factor único es ortogonal a todos los factores comunes y a todos los factores únicos asociados a otras variables.

Si se denota la existencia de una correlación entre cualesquiera dos variables, esta correlación es asumida como producto de factores comunes; en otras palabras, que la correlación parcial entre las dos variables, controlando por el factor común, dará por resultado 0.

Por la técnica del análisis factorial buscamos especificar un número hipotético mínimo de factores tal que todas las correlaciones parciales entre el resto de las variables devenga 0.

La varianza residual —es decir, la varianza que no se explica por los factores comunes— y la determinación de las comunales es uno de los problemas más complejos del análisis factorial.

Los diferentes métodos para la factorización están basados en distintos procedimientos para la estimación de las comunales, y el investigador debe estar consciente de las ventajas o desventajas de uno u otro para su diseño de investigación en particular.

Métodos de factorización en el SPSS

Existen disponibles 5 métodos de factorización: a) Factorización principal sin interacción (PA 1); b) Factorización principal con interacción (PA 2); c) Factorización canónica de Rao (RAO); d) Alfa factorización (ALPHA), y e) Imagen factorización (IMAGE).

Los cinco métodos tienen en común las siguientes características: 1) Todos los

factores son ortogonales; 2) Los factores son colocados en orden según su importancia; 3) El primer factor es comúnmente el factor general (es decir, tiene un factor de carga significativo en cada variable); el testeo de los factores tienden a ser bipolares (algunos factores de carga son positivos y otros negativos).

a) Factorización principal sin interacción

Se compone de dos métodos separados, según el usuario decida por: 1) Reemplazar la diagonal principal de la matriz de correlaciones por estimaciones de comunalidad; o 2) La diagonal principal de la matriz no se altera.

Que la diagonal se reemplace o no depende de que el investigador haya extraído sus factores iniciales, ya sea por el método de factores definidos o por el método de factores inferidos.

Cuando se utilizan estimaciones de comunalidad, estamos asumiendo la existencia de un factor único (U_j), y al reemplazar la diagonal estamos extrayendo los factores únicos de cada variable, analizando solamente las porciones remanentes de las mismas.

Cuando la diagonal de la matriz no se altera, los factores iniciales son definidos, consecuentemente los factores principales son calculados según los métodos especificados más arriba en la sección de factores definidos. En la matriz de componentes principales, los valores de peso asociados a cada componente representan la cantidad de varianza total que es explicada por el componente o factor. Esta estadística es calculada por el programa. En términos de fórmula uno puede representar entonces la varianza total y cada uno de los factores según la siguiente fórmula:

$$\sigma_T^2 = \sigma_a^2 + \sigma_b^2 + \dots + \sigma_n^2$$

Donde:

σ_T^2 = Varianza total;

σ_i^2 = Varianza explicada por el factor i .

A menos que haya indicación del usuario, el programa imprime y retiene únicamente componentes cuyo valor de peso sea igual o mayor que 1.0. El número de componentes significativos que van a ser retenidos para la rotación final son entonces determinados especificando un criterio mínimo de valor de peso.

b) Factorización principal con interacción

Es una modificación del primer método, sólo que aquí hay dos procedimientos que se hacen automáticamente: a) La diagonal principal es reemplazada con estimaciones de comunalidad; y b) Las estimaciones de comunalidad son corregidas por un proceso de

interacción en el que se van remplazando los elementos en la diagonal principal de manera tal que las diferencias entre dos comunales sucesivas sea negligible. Este método es el más recomendable para usuarios no familiarizados con los métodos de factorización.

c) Factorización canónica de Rao

Parte de los mismos supuestos del método clásico de factorización, y centra el problema alrededor de la estimación de la varianza única mediante una estimación de parámetros poblacionales a partir de datos muestrales.

Este tipo de factorización aplica un *test* de significación para el número de factores requeridos tal que la cantidad de factores requeridos por los datos y los factores hipotetizados no se desvíen significativamente del azar.

d) Alfa-factorización

Las *variables* incluidas en el modelo son considerados como una muestra del universo de *variables*. El propósito de esta factorización es definir los factores de manera tal que tengan un máximo de generalidad. Su aplicación tiene que ver con inferencia teórica más que estadística.

e) Imagen-factorización

Es un método bastante complicado desarrollado por Guttman para la determinación de comunales verdaderas. Por este método se obtiene una aproximación sobre la exacta proporción de la parte de la variable explicada por factores comunes y la parte de la variable explicada por el factor único.

C) Rotación

La rotación es un procedimiento por el cual se trata de encontrar una estructura tal que un vector aparezca como una función de un mínimo número de factores.

Este tercer último paso en la computación de un análisis factorial contiene diversas soluciones para la búsqueda de la mejor configuración, soluciones que dependen de los intereses teóricos y pragmáticos del investigador.

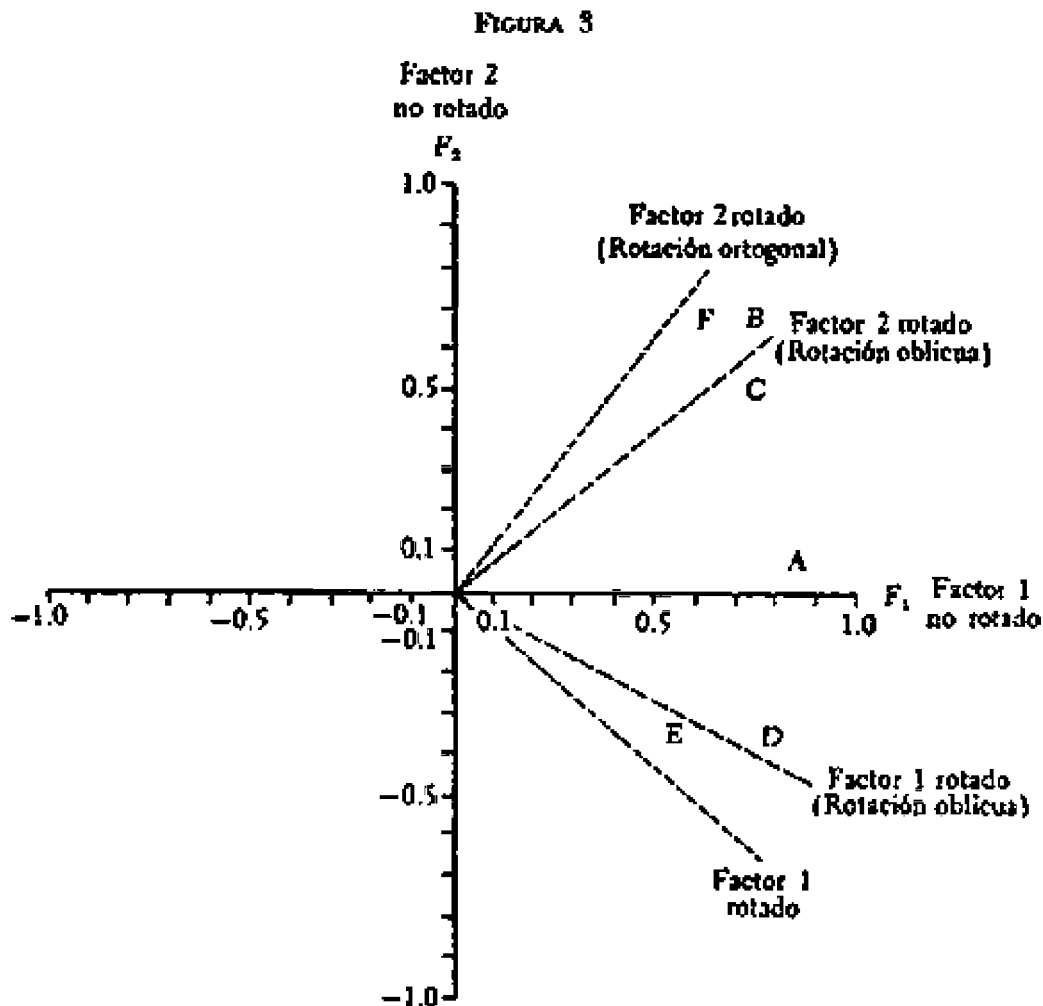
Para comenzar, éste debe decidirse por un método de rotación ortogonal o por un método de rotación oblicuo. Los métodos ortogonales proporcionan factores terminales no correlacionados, mientras que en los métodos oblicuos éstos pueden estar correlacionados. En términos gráficos una rotación ortogonal es aquella en la cual los ángulos entre los ejes se mantienen a 90°. En las rotaciones oblicuas los ángulos pueden ser agudos u obtusos.

La extracción de factores, tal como fue descrita en la sección de factorización, da lugar a una solución inicial que puede o no resultar en una estructura con significado.

Por medio de la rotación de los ejes y la solución de ecuaciones lineares simultáneas es posible interpretar de forma más adecuada la configuración de los resultados.

Una representación gráfica en la forma de un espacio de coordenadas cartesianas nos ayudará a entender en forma más clara el significado de las rotaciones. Vamos a representar solamente dos factores, ya que se trata de ejes cartesianos. Es posible representar espacialmente 3 factores (a través de planos), pero es imposible representar gráficamente más de tres.

F_1 y F_2 están representados por los ejes de referencia. Los factores de carga se representaban en el modelo por puntos (si el factor de carga en F_1 de una variable es .20 y en F_2 de .75, se avanza por el eje de las F_1 hasta la posición .20 y por el eje de las F_2 hasta .75; la línea de intersección señalará el punto para esa variable. Y así sucesivamente con las otras.



Los valores que figuran en la gráfica corresponden a la siguiente matriz con factores de carga no rotados.

<i>Variables</i>	<i>F₁</i>	<i>F₂</i>
A	.83	.10
B	.75	.63
C	.70	.50
D	.85	-.42
E	.60	-.40
F	.69	.57

La interpretación más elemental es que cuando más cerca están los puntos más relacionadas están las variables. Sin embargo, la interpretación a partir de estos dos factores es algo ambigua; por ejemplo, vemos que, si bien las variables están “cargadas” en el factor 1, el segundo factor es bipolar. Además, si bien la matriz principal provee información sobre factores de carga y comunales, no nos da una información precisa sobre la estructura de las relaciones. De por sí la matriz original (no rotada) es arbitraria, en el sentido de que se pueden trazar infinito número de posiciones. La rotación de los ejes es entonces un medio para buscar la mejor manera de acomodar los datos en un espacio n-dimensional.

Una razón adicional para la rotación, como señala Kim (*op. cit.*), es que los factores de carga en la solución no rotada dependen muy estrechamente en el número de variables. Los factores rotados son más estables. En fin, el objetivo de la rotación es la obtención de factores teóricamente significativos.

Las líneas de puntos en la gráfica señalan una rotación ortogonal para los datos que aparecen en el cuadro anterior. Las líneas segmentales señalan una rotación oblicua.

Por lo general, los cuadros rotados proporcionan patrones mucho más claros. Siguiendo el ejemplo hipotético presentado, los factores rotados y no rotados serían ahora (rotación ortogonal):

<i>Variables</i>	<i>Factores no rotados</i>		<i>Factores rotados</i>	
	<i>F₁</i>	<i>F₂</i>	<i>F₁</i>	<i>F₂</i>
A	.83	.10	.60	.58
B	.75	.63	.14	.96
C	.70	.50	.22	.83
D	.85	-.42	.94	.09
E	.60	-.40	.72	.08
F	.69	.57	.14	.90

La estructura de los datos es ahora mucho más clara, aunque la variable *A* sigue

teniendo carga en los dos factores. Por lo demás queda claro que en el factor 1 intervienen las variables *D* y *E*, mientras que en el factor 2 la carga se concentra en las variables *D*, *C* y *F*.

Las rotaciones, cualquiera que sea su tipo, tienen como objeto hacer que los valores en la horizontal o en la vertical de los ejes se aproximen lo máximo posible a 0. Al hacer esto con un factor, maximizamos a la vez el valor del otro factor.

Los métodos de rotación incluidos en el SPSS son: i) Rotación ortogonal *quartimax*; ii) Rotación ortogonal *varimax*; iii) Rotación ortogonal *equimax*; iv) Rotación oblicua.

i) *Rotación ortogonal quartimax*. Esta rotación sigue el principio de reducción del máximo de complejidad en una variable, mediante la rotación de los factores iniciales, de tal manera que el factor de carga se concentre en un factor, haciendo que el peso de los otros factores se acerque lo máximo al valor 0. Este método destaca la simplificación de las líneas, por lo tanto, el primer factor rotado tiende a ser un factor general (muchas variables tienden a concentrar su peso en él). Los siguientes factores tienden a ser *subconglomerados* de variables.

ii) *Rotación ortogonal varimax*. El método *varimax* se concentra en la simplificación de las columnas de la matriz inicial. Es el método de uso más generalizado.

iii) *Rotación ortogonal equimax*. Es un método intermedio a los dos anteriores. En vez de concentrarse en la simplificación de las líneas (*quartimax*), o en la simplificación de las columnas (*varimax*), *equimax* trata de lograr algo de cada una de esas simplificaciones.

iv) *Rotación oblicua*. En las rotaciones oblicuas se acepta por principio que los factores están intercorrelacionados. A partir de esto, los ejes son rotados libremente de manera tal que los hiperplanos se coloquen oblicuos los unos a los otros. Hay diversos métodos de rotación oblicua (a partir de los gráficos de las rotaciones ortogonales, el método del plano único de Thurstone, de rotación directa hacia estructuras primarias de Harris, etc.). El método utilizado en el subprograma es un método de rotación oblicua objetivo, utiliza un método directo para la simplificación de los factores primarios de carga. Los factores son intercorrelacionados, si tales intercorrelaciones existen; sin embargo, algunos métodos tienden a hacer que los factores resultantes estén más correlacionados que otros métodos. En el método rotacional oblicuo, los valores δ son colocados en 0 (son los que por lo general tienden a reproducir las mejores soluciones oblicuas). Sin embargo, el usuario puede alterar los valores δ para lograr menor o mayor oblicuidad. Cuando varios valores δ son especificados, el programa calculará una rotación oblicua para cada valor especificado.

D) Interpretación

Es la tarea teórica de identificar el contenido y la naturaleza de los factores. Esto se hace mediante procesos inferenciales acerca de qué tienen en común las variables con alta carga con las variables de carga moderada, o con las variables con factores de carga próximos a cero. Esas inferencias son probadas *a posteriori* en otros diseños con las

hipótesis necesarias, etcétera.

Producto del programa factor

Un cuadro de análisis factorial completo provee la siguiente información:

1) *Una matriz de correlaciones de las variables* en el modelo tal como aparece en la figura.

2) *Factores de carga iniciales*. Este cuadro contiene factores ortogonales proporcionando el patrón y la estructura de la matriz. Ya sea que los factores sean inferidos o definidos, se ordenan por importancia decreciente.

Cuando existe un solo factor se dice que el conjunto de variables es puro, o saturado, o cargado con el factor; cuando hay más de un factor, se dice que el conjunto es factorialmente complejo. Este cuadro nos informa tanto sobre el número de factores como de la magnitud de la carga o saturación de cada variable en cada uno de los factores iniciales. Los factores de carga varían de -1.0 , pasando por 0 , hasta $+1.0$ y se interpretan de la misma manera que un coeficiente de correlación (de hecho los factores de carga expresan la correlación entre las distintas variables y los factores). Las comunalidades son la suma de los cuadrados de los factores de carga en una variable y expresan el factor de varianza común: $h^2 = (F_1^2) + (F_2^2) + \dots (F_n^2)$. Con base en la matriz inicial, el investigador puede decidir sobre la cantidad final de factores a retener.

3) *Pesos para estimar variables a partir de factores (factor-pattern matrix)*. Contiene los pesos de regresiones de los factores comunes y nos informa sobre la composición de una variable en términos de factores hipotéticos. Esta matriz es rotada y nos permite expresar la variable como una combinación de variables independientes, sean éstas definidas o inferidas.

4) *Pesos para estimar factores a partir de variables (factor-estimate o factor-score matrix)*. Provee medios para estimar puntajes en los factores a partir de las variables observadas.

5) *Correlación entre factores y variables (factor-structure matrix)*. Proporciona el coeficiente de correlación entre cada variable y cada factor. La solución es rotada, siendo este cuadro idéntico a la *factor-pattern matrix*.

6) *Matriz de correlación para los factores terminales*. La interpretación de los cuadros es diferente, según la solución haya sido ortogonal u oblicua. Las matrices básicas para los dos tipos de soluciones son:

	Solución ortogonal	Solución oblicua
Datos básicos	Matriz de correlación idéntica	
Factores iniciales	Idéntica matriz factorial ortogonal	
Factores terminales	Factor-matrix	a) Pattern-matrix b) Structure-matrix
Correlación entre factores		Factor-correlation matrix
Factor-estimate matrix	Factor-estimate matrix	Factor-estimate matrix

Soluciones terminales para factores rotados ortogonalmente

Los coeficientes en el cuadro representan tanto pesos de regresión (*pattern-matrix*) como coeficientes de correlación (*structure-matrix*). Esto es porque la solución es ortogonal. El ejemplo de Kim nos ayudará a interpretar el cuadro:

Matriz factorial final con rotación varimax

Variables	Factor 1	Factor 2	Factor 3
A	.88920	.07629	.03290
B	.78523	.11023	.03768
C	.10210	.67532	.06342
D	.07297	.85632	.09643
E	.08390	.09170	.76480
F	.12515	.00320	.69532
G	.32460	.31210	.01271

Examinando cada variable (es decir, cada línea), vemos que el determinante más importante en la variable *A* es el factor 1, lo mismo que en la variable *B*. Para las variables *C* y *D* el determinante más importante es el factor 2. Para las variables *E* y *F* el factor más importante es el 3. Todas estas variables tienen por consiguiente una complejidad de 1. La variable *G*, por el contrario, tiene como determinantes principales los factores 1 y 2, por lo tanto su complejidad factorial es de 2, es decir que ésta no es una variable que se explique por una dimensión, sino que mide dos dimensiones; y así sucesivamente para complejidades factoriales 3, 4 ... *n*.

Leyendo ahora las columnas, podemos determinar cómo los factores hipotéticos (cada uno de los factores que aparecen en el cuadro y que representan las variables independientes que explican cada una de las variables en el cuadro) dan cuenta de cierta proporción de la varianza en la variable dependiente.

Por ejemplo, la varianza de la variable *A* explicada por el factor 1 es $(.8892)^2 = .79067$, es decir, 79% de la varianza de *A* es explicada por el factor 1.

Siguiendo el mismo razonamiento, la proporción de la varianza explicada por todos los factores, en el caso de la variable A , es:

$$h_A^2 = (.88920)^2 + (.07829)^2 + (.03230)^2 = .79783$$

Esto es lo que llamamos más arriba comunalidad, y es claro que en el caso de la variable A la contribución de los factores 2 y 3 es mínima (exactamente .00716).

El complemento de la comunidad ($1 - h_j^2$) representa la proporción de la varianza única, es decir, la proporción de la varianza no explicada por los factores comunes, o por ninguna variable en el conjunto.

Con los datos de la matriz, el investigador también puede calcular los coeficientes de correlación entre cualquier par de líneas, determinando así las fuentes de variación común a ambas variables. Por ejemplo, la correlación entre la variable A y la variable B en los tres factores es:

$$R_{12} = r_{1F_1}r_{2F_1} + r_{1F_2}r_{2F_2} + r_{1F_3}r_{2F_3} = (.88920)(.78523) + (.07829)(.14023) + (.03230)(.05768) = .71105$$

La correlación entre A y B es debida básicamente al factor 1. Mediante la fórmula se pueden calcular las intercorrelaciones para todos los pares de variables.

Soluciones terminales para factores rotados oblicuamente

En una solución oblicua habrá dos matrices separadas, una para la *pattern-matrix* y otra para la *structure-matrix*. La *pattern-matrix* delinea más claramente la agrupación de variables. El cuadrado de un *pattern*-coeficiente representa la contribución directa de un factor determinado a la varianza de una variable. En la medida en que un factor puede contribuir a la varianza de una variable a través de otros factores correlacionados (es decir, una contribución indirecta), el total de la varianza del que da cuenta un factor no es igual a la suma de las contribuciones directas.

La matriz de estructuras (*structure-matrix*) consiste en coeficientes de correlación. La contribución entre la variable A y el factor 1 en el ejemplo que sigue es .97652 y su elevación al cuadrado dará cuenta de la cantidad de varianza en la variable A explicada por el factor 1.

Matrices para factores oblicuos

Variables	Patterns-matrix		Structure-matrix	
	Factor 1	Factor 2	Factor 1	Factor 2
A	.99978	-.11012	.97652	.08240
B	.88724	-.08012	.87234	.08982
C	.76098	.34980	.82600	.43000
D	-.08202	.99870	.11721	.99668
E	.05368	.96728	.24231	.97823

El subprograma *factor* provee una *representación gráfica* de los tres métodos de rotaciones ortogonales. Ya que la representación se hace en ejes cartesianos, los factores son tomados de dos en dos, es decir, si hay tres factores habrá tres gráficos (F_1 con F_2 ; F_1 con F_3 y F_2 con F_3). Por este gráfico, como ya señalamos más arriba, el usuario tiene una idea más clara sobre el agrupamiento de variables, sus valores, etc. Pudiendo utilizarlos para decisiones sobre la rotación oblicua por ejemplo.

El programa *factor* también construye índices compuestos que representan las dimensiones teóricas asociadas a los respectivos factores.

ANÁLISIS DISCRIMINANTE (SUBPROGRAM DISCRIMINANT)

Es una técnica estadística para la clasificación, predicción y análisis en problemas de grupos o clases de objetos. Puede ser utilizado tanto para la determinación de las diferencias entre dos o más grupos como para —a partir de esas diferencias— construir esquemas clasificatorios de manera que sea posible clasificar cualquier caso cuya pertenencia a un grupo específico nos es desconocida.

Los supuestos que subyacen en el análisis discriminante son: *a)* Los grupos son discretos e identificables; *b)* Cada observación en los grupos puede ser descrita por un conjunto de mediciones de m características o variables; *c)* Las m variables tienen una distribución normal multivariata en cada población.

La distinción entre los grupos se realiza a partir de un conjunto de variables discriminatorias, esto es, variables que el investigador sospecha que miden características sobre las cuales los m grupos difieren. Por medio del análisis discriminante esas variables son combinadas linealmente de manera que se maximice la distinción entre los grupos.

Las combinaciones entre las funciones discriminatorias toman la forma de una ecuación donde las *funciones discriminantes* son:

$$D_i = d_{i1}Z_1 + d_{i2}Z_2 + \dots + d_{im}Z_m$$

Donde:

D_i = Puntaje de la función discriminante i .

d_i = Coeficiente de carga.

Z_i = Valor estándar de las m variables discriminantes utilizadas.

Una vez que se determinan las funciones discriminantes, éstas pueden ser utilizadas para propósitos de clasificación y de análisis.

Como técnica *clasificatoria*, el análisis discriminante puede ser utilizado para clasificar cualquier caso cuya pertenencia a un grupo específico nos es desconocida. Es decir, se construyen un conjunto de reglas para clasificar las observaciones en el grupo más apropiado. Por ejemplo, una vez determinadas las características que diferencian a conservadores de radicales, podemos utilizar las combinaciones lineales de las variables discriminatorias en la determinación de la probabilidad de que un miembro cualquiera cuya afiliación desconocemos se adhiera a conservadores o radicales.

Como *técnica analítica* el análisis discriminante permite detectar en qué medida las variables discriminatorias efectivamente discriminan cuando se combinan en funciones discriminantes. También es posible reducir el número de funciones discriminantes, siguiendo el mismo tipo de razonamiento utilizado en el análisis factorial. Por medio de la técnica es posible el estudio de las relaciones espaciales entre los grupos, así como identificar las variables que contribuyen de manera más significativa a la diferenciación entre los m grupos.

Los pasos en el análisis discriminante son los siguientes:

1) *El investigador selecciona un conjunto de variables que sospecha que van a diferenciar entre los m grupos. Las variables pueden ser tantas como el investigador considere necesarias.*

2) *Las variables seleccionadas en el primer paso pueden a su vez ser seleccionadas o no para su inclusión en el análisis discriminante. Si el investigador decide incluir todas sus variables, se deben dar instrucciones (rutina *Method = direct*) por medio de las cuales las funciones discriminantes se crearán independientemente del poder discriminatorio de cada una de las variables independientes.*

Si, por el contrario, el investigador decide seleccionar sus variables en base al poder discriminatorio, es decir, si decide introducir un criterio estadístico adicional a los teóricos, existen 5 criterios de selección disponibles en el SPSS. Los procedimientos de selección operan de manera que seleccionan primero la variable que tiene el valor más alto en el criterio de selección; luego esa variable es apareada con cada una de las otras variables hasta seleccionar una segunda variable que, combinada con la primera, mejora el criterio de selección; luego se aparean estas dos variables con cada una de las que quedan hasta seleccionar una tercera variable que combinada con las dos anteriores mejora aún más el criterio; y así sucesivamente, hasta que la inclusión de una variable adicional no provea un mejoramiento en la discriminación entre los grupos.

Estos métodos de ubicación de las variables por rango y de evaluación de su poder discriminatorio no tienen la precisión de los análisis de regresión, ya que no existe una prueba clara para determinar la significancia de un coeficiente en particular en una

función discriminativa dada. De ahí que los 5 métodos que se mencionan subrayan distintos aspectos de la separación.

Método de Wilks (Method = Wilks). Está basado en una prueba de significación de las diferencias entre grupos, mediante el *test F* de las diferencias entre grupos centroides. El método toma en consideración la diferencia entre todos los centroides y la cohesión entre los grupos. Las variables pueden ser ordenadas según los valores de los coeficientes λ de Wilks de manera que se dé un rangueamiento en referencia a su poder discriminatorio relativo. A más bajo valor de λ , más alto poder discriminatorio de la variable. El inconveniente de este método es que deja de considerar las correlaciones entre el conjunto de variables que se está utilizando, de manera que solamente en los casos en que las variables sean independientes (no correlacionadas) este método permitirá un ordenamiento y comparación válidos.

El *método Mahal (Method = Mahal)* maximiza las distancias mahalanobis entre los dos grupos más próximos.

El *método Miniresid (Method = Miniresid)* separa los grupos de manera que la variación residual sea mínima. Tomando en cuenta la correlación múltiple entre el conjunto de variables discriminantes y una variable muda que identifica el par de grupos correspondiente, su objetivo es minimizar R , esto es, la variación residual.

El *método Maxminf (Method = Maxminf)* maximiza la distancia entre los grupos de manera que se seleccione la razón R más pequeña entre dos pares de grupos más próximos.

El *método de Rao (Method = Rao)* utiliza una medida de distancia V . La variable seleccionada es aquella que contribuye al aumento más grande en V cuando se agrega a las otras variables, de manera que se obtenga una separación máxima entre los grupos. Estas decisiones tienen que ver con la comparación del poder discriminatorio en diferentes conjuntos de variables (diferentes en relación a su tamaño). La solución de Rao tiene que ver con la significancia que puede tener la agregación o no de una variable en particular.

3) *Determinación del número de funciones discriminantes*. Para cualquier cantidad de grupos o cualquier cantidad de variables discriminantes, el máximo número a derivar será igual a la cantidad de variables discriminantes o a la cantidad de grupos menos 1, cualquiera que sea menor. Es decir que tres funciones pueden ser suficientes para describir a dos grupos. El subprograma *discriminant* provee dos medidas para juzgar la importancia de las funciones discriminantes: porcentaje relativo del *eigenvalue* asociado a la función y la significación estadística de la información discriminatoria.

El *eigenvalue* es una medida especial que representa los valores característicos de una matriz cuadrada, y es una medida relativa de la importancia de la función, en la medida en que la suma de los *eigenvalues* es una medida de la variancia total que existe entre las funciones discriminantes; así un *eigenvalue* en particular es expresado como un porcentaje de esa suma. En la medida en que las funciones discriminantes son derivadas por orden de importancia, el proceso de derivación es para cuando el porcentaje relativo (es decir, el valor del *eigenvalue*) es demasiado pequeño.

El segundo criterio para juzgar la importancia de las funciones discriminantes es el *test* de la significación estadística de la información discriminante todavía no contemplada por las funciones ya determinadas. El método de cálculo utilizado es el λ de Wilks.

4) *Interpretación de los coeficientes de función discriminante.* Habíamos visto más arriba que la ecuación de las funciones discriminantes era:

$$D_i = d_{i1}Z_1 + d_{i2}Z_2 + \dots + d_{im}Z_m$$

Los coeficientes de la función discriminante corresponden a los d_{ij} y son utilizados para computar el puntaje discriminante que es el resultado de aplicar la fórmula que aparece más arriba. Habrá por lo tanto un puntaje separado para cada caso en cada función. En la medida en que los puntajes Z son estándar, su media es 0 y su valor estándar es 1. Por lo tanto, cualquier puntaje singular representa una desviación de la media de todos los casos sobre una función discriminante dada.

Computando el promedio de los puntajes para un grupo en particular, tenemos calculada la media del grupo en la función discriminante respectiva. Para cada grupo, las medias de todas las funciones se denominan grupo centroide; éste señala la localización más típica de los casos en ese grupo con referencia al espacio de la función discriminante. Una comparación de las medias de los distintos grupos en cada función nos indica entonces cómo se distribuyen los grupos a lo largo de una dimensión.

Los coeficientes estandarizados de las funciones discriminantes se pueden interpretar en forma similar a los coeficientes beta en las regresiones múltiples. El signo del coeficiente nos indica si la contribución de la variable es positiva o negativa. El valor del coeficiente indica la contribución relativa de la variable a la función. Los coeficientes estandarizados pueden ser utilizados —como en el análisis factorial— para identificar las características dominantes que ellos miden, nombrándolos así según características teóricas.

El programa contiene una opción para el cálculo de *coeficientes no estandarizados*, útiles para propósitos computacionales, pero que no nos informan sobre la importancia relativa de las variables.

5) *Distribución gráfica (plots) de los puntajes discriminantes.* El programa imprime una representación espacial de la distribución de los puntajes discriminantes a lo largo del continuo de dos primeras funciones discriminantes. Es posible obtener ya sea la distribución de todos los casos en un gráfico, o gráficos separados para cada grupo. La representación espacial es particularmente útil para el estudio de los grupos centroides y su localización relativa, así como también para un análisis del grado en que los grupos se superimponen. Cuando solamente existe una función discriminante, la distribución toma la forma de un histograma.

6) *Rotación de los ejes de las funciones discriminantes.* Tal como en el análisis factorial, es posible rotar la orientación espacial de los ejes manteniendo constante la localización relativa de los casos y de los centroides. Un criterio puede ser la solución

varimax, esto es, hacer rotar los ejes de manera que las variables discriminantes se aproximen a 1.0 o a 0.0. Mediante esto, si bien es posible mejorar la interpretación de la distribución de las variables principales, hay pérdida de información referente a la importancia relativa de cada función.

7) *Clasificación de casos*. Por este proceso es posible identificar la pertenencia de un caso a un grupo determinado, cuando solamente conocemos los valores del caso en las variables discriminatorias. La clasificación se logra mediante el uso de una serie de funciones de clasificación, una para cada grupo. Se computan para cada caso tantos puntajes como grupos existen y el caso es clasificado en el grupo con el puntaje más alto. Los puntajes pueden ser asimismo en probabilidades de pertenencia a grupo asignándose el caso al grupo cuya probabilidad de pertenencia es más alto.

El sistema de clasificaciones de probabilidad es útil no solamente para adjudicar casos en un grupo, sino además para controlar cuán efectivas son las variables discriminantes. De ahí que sus valores se utilicen aun para los casos de selección de variables y de funciones. Si existe un número de casos cuya afiliación conocemos, pero que están mal clasificados, entonces las variables seleccionadas son muy pobres en el proceso de discriminación.

Algunos ejemplos de análisis discriminante

a) Ejemplo en dos grupos

El ejemplo es de una aplicación de Heyck y Klecka, y aparece en el *Manual* del SPSS. Se trata de un análisis del Parlamento Británico durante el periodo 1874-1895, en el que el Partido Liberal estaba fraccionado entre radicales y no radicales. Algunos miembros del Parlamento fueron clasificados en uno u otro grupo según documentos históricos. Sin embargo, quedaron sin clasificar un conjunto de miembros para los cuales no se disponía de suficiente información o para los cuales la información era contradictoria.

Las variables discriminantes seleccionadas fueron votos en el Parlamento, los votos en asuntos particularmente relacionados al programa radical. Se seleccionaron 17 de estas variables, siendo los asuntos a votar los siguientes:

Fecha

25/marzo /74	1. Horas de votación
17/abril /74	2. Gastos de la Corona
10/junio /74	3. Educación no sectaria
17/junio /74	4. <i>Temperance reform</i>
9/junio /75	5. Escolarización obligatoria
17/junio /75	6. Parlamentos trienales
14/junio /75	7. Asignaciones de tierras
15/julio /75	8. Gastos de la Corona
15/julio /75	9. Extensión
5/abril /76	10. Educación gratuita
30/mayo /76	11. Extensión
10/julio /76	12. Control del estado en educación
13/marzo /77	13. <i>Temperance reform</i>
23/marzo /77	14. Reformas en Turquía
13/mayo /80	15. Prerrogativas de la Corona en política externa
24/febrero/80	16. Parlamentos quinquenales
5/marzo /80	17. <i>Temperance reform</i>

Los votos desde el punto de vista radical fueron clasificados con +1, los votos en contra con —1 y las abstenciones con 0.

Antes de calcular el análisis discriminante las variables fueron seleccionadas para su inclusión mediante el método *V* de Rao, así solamente se seleccionaron 11 de las 17 ocasiones de voto. Esos 11 votos proveían un alto grado de separación entre los grupos (el lambda de Wilks dio como valor .19264 con una correlación canónica de .899 para la función discriminante). Las variables eliminadas fueron la 6, 7, 8, 9, 10 y 17.

Los valores obtenidos después del análisis discriminante para las 11 variables discriminatorias fueron los siguientes:

<i>Paso número</i>	<i>Variable</i>	<i>F para incluir o extraer del análisis</i>	<i>Número incluido</i>	<i>Lambda de Wilks</i>	<i>Rao V</i>	<i>Cambio en el V de Rao</i>
1	Var018	39.95802	1	0.06125	39.97799	39.95799
2	Var011	33.86707	2	0.45926	91.83968	51.88168
3	Var014	14.96441	3	0.38370	125.28114	33.44147
4	Var005	10.68981	4	0.33584	164.25507	23.97392
5	Var015	8.99831	5	0.29918	182.49631	23.24174
6	Var003	8.65288	6	0.26770	213.37398	30.87717
7	Var002	6.10148	7	0.24678	238.06580	24.69182
8	Var001	5.55072	8	0.22889	262.77515	24.70935
9	Var016	6.10082	9	0.21054	292.47461	29.69946
10	Var012	3.86319	10	0.20070	310.63940	18.16479
11	Var004	2.84350	11	0.19264	326.89038	16.25098

Clasificación de coeficientes de función

	<i>Grupo radical</i>	<i>Grupo no radical</i>	<i>Coefficientes discriminantes</i>
Var001	4.78972	1.72740	-0.340
Var002	3.77127	0.07671	-0.410
Var003	3.41047	1.11915	-0.254
Var004	0.40823	-1.01541	-0.158
Var005	5.96921	2.52084	-0.382
Var011	3.94122	1.46606	-0.275
Var012	3.61321	0.97632	-0.293
Var013	3.12955	0.53350	-0.268
Var014	1.31335	-0.79479	-0.234
Var015	1.21796	-0.86230	-0.231
Var016	2.82108	-0.42259	-0.360
Constant	-12.65345	-1.92266	1.167

<i>Funciones discr minantes</i>	<i>Eigen- valor</i>	<i>Porcen- taje re- lativo</i>	<i>Correla- ción ca- nónica</i>	<i>Funcio- nes de- rivadas</i>	<i>Lambda de Wilks</i>	<i>ji cuadrado</i>	<i>DF</i>
1	4.19092	100.00	0.899	0	0.1926	119.401	11

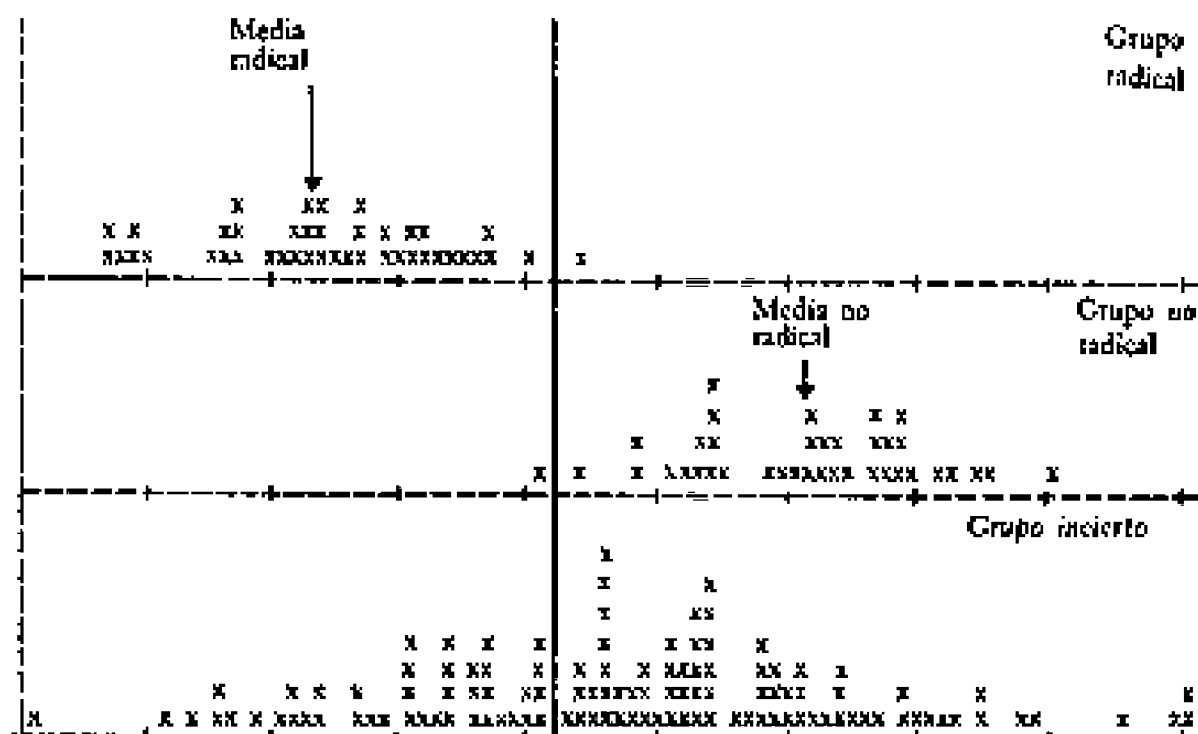
Coefficientes estandarizados de función discriminante

Var001	—0.17241
Var002	—0.20581
Var003	—0.15082
Var004	—0.10021
Var005	—0.16783
Var011	—0.18131
Var012	—0.13519
Var013	—0.20311
Var014	—0.10893
Var015	—0.16580
Var016	—0.17411

Centroides de grupo en espacio reducido

Grupo 1	—0.83134
Grupo 2	0.91885

FIGURA 4. Distribución gráfica (plots) de los puntajes discriminantes para miembros del Parlamento Inglés; 1874-1895



Los coeficientes discriminantes no estandarizados son utilizados para calcular el puntaje discriminante para la función. Esto se obtiene multiplicando cada coeficiente por el valor en la variable (voto) y luego sumando los productos más la constante. Supongamos los votos de dos miembros del parlamento, sir Wilfrid Lawson y lord Hartington.

En la votación de los diferentes asuntos los miembros lo hicieron de la siguiente forma (+1 indica voto radical, —1 en contra y 0 abstención).

<i>Variable</i>	<i>Lawson</i>	<i>Hartington</i>
1	+1	0
2	0	0
3	+1	+1
4	+1	0
5	+1	0
11	+1	+1
12	+1	0
13	+1	-1
14	+1	0
15	+1	-1
16	+1	0

Los puntajes para los dos miembros serán entonces:

Lawson

$$D = (-.340) (+1) + (-.410) (0) + (-.254) (+1) + (-.158) (+1) + (-.382) (+1) + (-.275) (+1) + (-.293) (+1) + (-.288) (+1) + (-.234) (+1) + (-.231) (+1) + (-.360) (+1) + (1.167) = -1.648$$

Hartington

$$D = (-.340) (0) + (-.410) (0) + (-.254) (+1) + (-.158) (0) + (-.382) (0) + (-.275) (+1) + (-.293) (0) + (-.288) (-1) + (-.234) (0) + (-.231) (-1) + (-.360) (-1) + (1.157) = 1.157$$

Los puntajes calculados para todos los miembros aparecen en la distribución gráfica. Donde 0 es la media total de todos los radicales y no-radicales conocidos.

El siguiente paso es la clasificación de los miembros no identificados claramente como radicales o no-radicales. Esto se hace calculando para cada uno de ellos el puntaje discriminante, usando los coeficientes derivados de los radicales y no-radicales conocidos. La lógica de la clasificación se basa en la comparación de las pautas de votos de radicales conocidos con las pautas de votos del miembro incierto, de manera de clasificarlo en el grupo más similar en cuanto a pauta. Para ello se utilizan una serie de ecuaciones (una para cada grupo), cada ecuación dará lugar a una probabilidad, y el

miembro será colocado en el grupo para el cual obtiene una probabilidad más grande. Este sistema tiene algunos problemas, sobre todo cuando las probabilidades de pertenencia a uno u otro grupo son del tipo .53, .47. Sucede en el caso del ejemplo que consideramos que algunos miembros no muestran una pauta consistente de apoyo a uno u otro grupo.

b) Ejemplo con varios grupos

El ejemplo es de Gansner, Seegrist y Walton ⁶ en el que se trata de seleccionar subregiones en el Estado de Pennsylvania (Estados Unidos) mediante una combinación de análisis discriminante y técnicas de agrupamiento (*Clustering techniques*). La regionalización está relacionada con un análisis de la eficiencia económica de la industria de la madera, con el objetivo de desarrollar un sistema más eficiente de tala y entrega de derivados de la madera. Se busca que las subregiones sean homogéneas y compactas en términos de producción de madera y de mercado potencial. Se busca además que cada subregión tenga fronteras equidistantes de sus centros de consumo.

Las variables discriminantes seleccionadas fueron las siguientes:

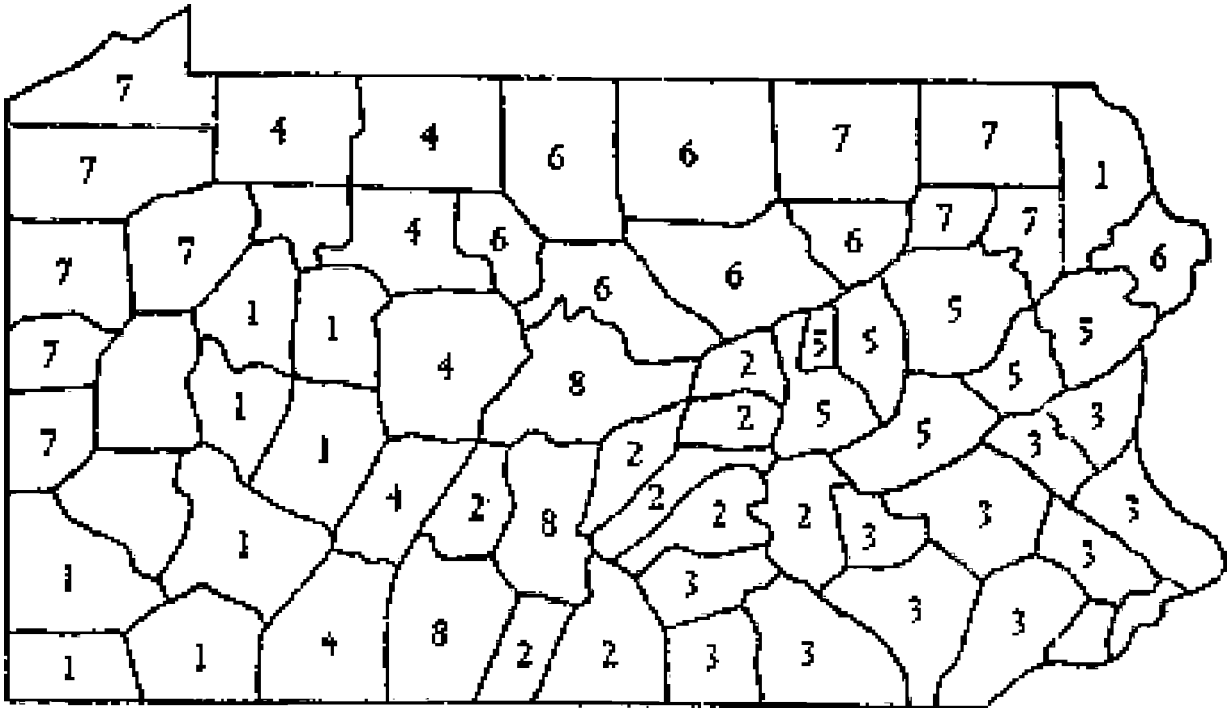
<i>Variables que reflejen existencia de madera y aceptabilidad</i>	<i>Variables que reflejen actividad de mercado y demanda</i>
1: Bosques comercializables como % del total de tierras.	9: Promedio anual de producción de pulpa de madera (<i>round pulpwood</i>).
2: Volumen de <i>stock</i> creciente por acre de bosque comercializable.	10: Madera blanda (<i>softwood</i>) como % de la producción de pulpa de madera.
3: Volumen de <i>stock</i> creciente en bosques comercializables.	11: Número de fábricas que usan pulpa de madera dentro de un radio de 100 millas de radio del centro del municipio.
4: Madera blanda (<i>softwood</i>) como % del total del volumen de <i>stock</i> creciente.	12: Capacidad total productiva de las fábricas que usan pulpa de madera.
5: % de bosques en propiedad pública.	13: Número de campos madereros y firmas contratantes.
6: Diferencia de elevación máxima.	14: Empleo en campos madereros y firmas contratantes.
7: % de superficie con pendiente suave.	15: Número total de industrias madereras y de productos derivados de la madera.
8: Millas de carreteras por milla cuadrada de superficie total.	

Los grupos son todos los municipios de Pensilvania, excepto Delaware y Filadelfia,

esto es, 65 municipios.

El segundo paso utilizado por los autores fue el agrupamiento de municipios. La técnica utilizada fue el análisis discriminante *stepwise* que corresponde a nuestro segundo paso (selección de variables para su inclusión en el análisis). Resultaron seleccionados 8 grupos que forman subregiones compactas, tal como aparecen en la figura 5:

FIGURA 5. Ocho grupos de municipios en Pensilvania formados por análisis discriminantes.



A continuación se calcularon las medidas de distancia generalizada (D^2) o medida de similitud entre pares de grupos; basada en valores F , cuyos resultados fueron los siguientes:

Tabla 2. Valores D^2 .

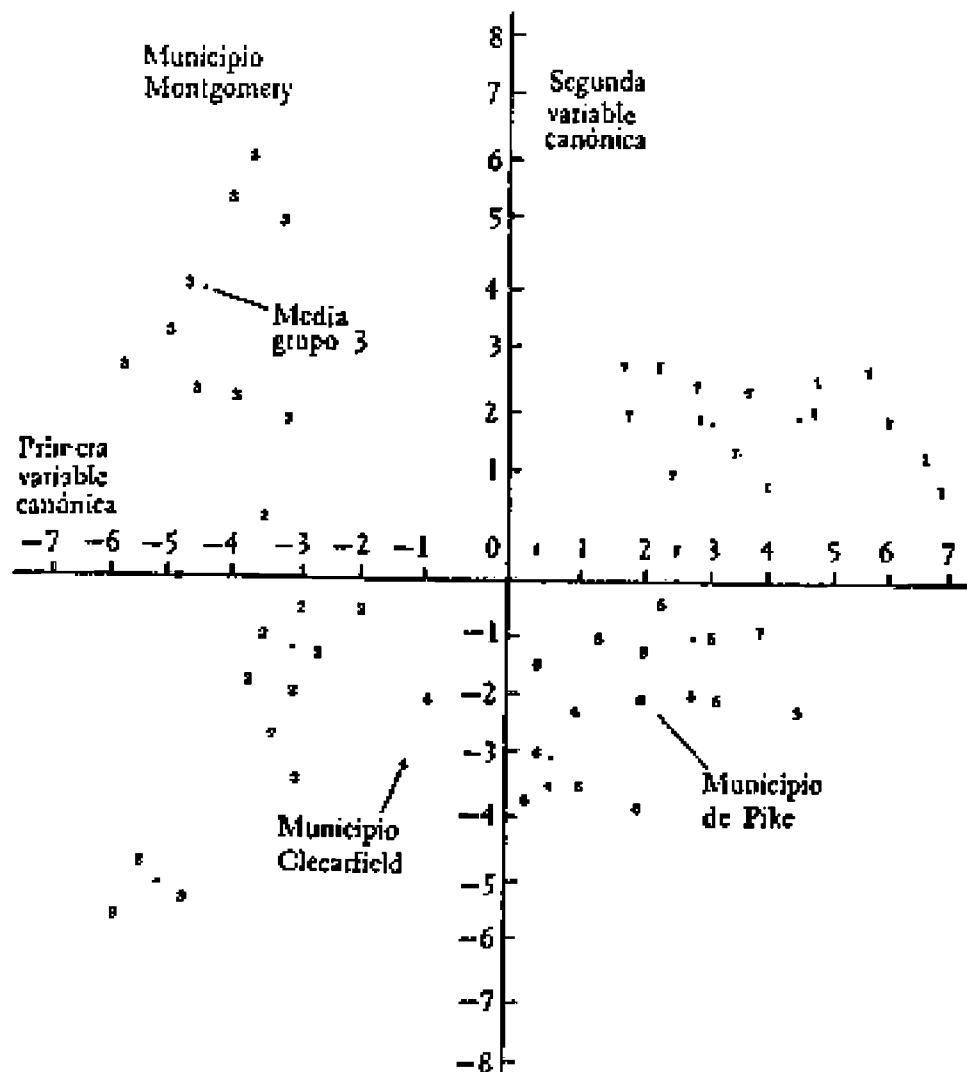
Grupo	1	2	3	4	5	6	7
1	205.2	...	...	...	...	...	...
2	1279.6	...	...	...	...	...	...
3	1864.3	287.4	...	...	...	...	...
4	506.1	494.2	1219.3	...	...	...	...
5	785.3	703.1	2334.3	1497.7	...	...	...
6	600.9	161.2	1335.8	251.0	664.5	...	...
7	158.9	783.9	1112.2	492.8	396.8	394.8	...
8	2476.4	210.4	1204.6	398.3	1612.5	627.8	1698.5

Como puede verse en la tabla, los grupos 1 y 7 son muy similares (tienen los valores D^2 más bajos). Los grupos 2 y 3; 2 y 6; 2 y 8; 4 y 6 también son similares. Por el contrario, los grupos 1 y 3; 1 y 8; 3 y 5; 5 y 8; 7 y 8 son muy diferentes.

El programa suministra una distribución gráfica de los valores de las dos primeras variables canónicas que muestra gráficamente la distribución en los grupos. Las medias de los grupos aparecen identificadas con asteriscos.

Los autores no proporcionan todos los valores, pero en la figura 6 puede verse que la media del grupo 3 está localizada a -4.211 en la primera variable canónica y a $+4.016$ en la segunda. La municipalidad de Montgomery que pertenece al grupo 3 tiene como coordenadas $(-3.237$ y $+5.959)$. Las variables canónicas están relacionadas con los valores D^2 y reflejando la similaridad entre los grupos. Puede verse en la figura 6 que los grupos 1 y 7 están próximos, mientras que 1 y 8 están alejados.

FIGURA 6. Municipios agrupados en 8 grupos



El tercer círculo realizado se refiere a una matriz de clasificación para evaluar las probabilidades de clasificar los municipios en grupos. El cuadro 3 muestra una matriz casi perfecta, con la excepción de tres elementos, lo que indica que la clasificación es casi perfecta.

Con base en la similaridad de valores D^2 se forman ahora 5 subregiones:

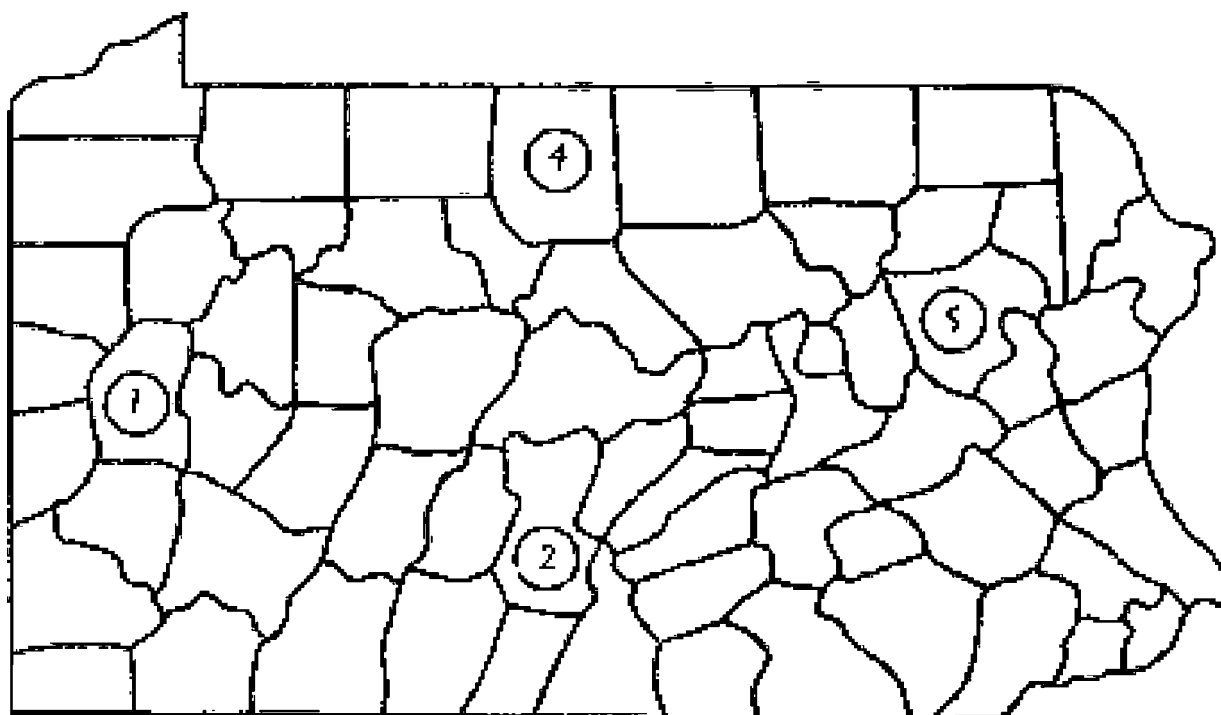
CUADRO 3. Matriz para la combinación de asignación de municipios en 8 grupos

Subregión	Número de municipios por subregiones							
	1	2	3	4	5	6	7	8
1	9	0	0	0	0	0	1	0
2	0	9	11	0	0	11	0	0
3	0	0	11	0	0	0	0	0
4	0	11	0	6	0	0	0	0
5	0	0	0	0	7	0	0	0
6	0	11	0	0	0	7	1	0
7	1	0	0	0	0	0	10	0
8	11	0	0	0	0	0	0	3

La combinación de los grupos 1 y el noroeste de la región 7 produce una unidad. Los grupos 2 y 8, otra unidad compacta. El grupo 3 queda como está, se combina el grupo 5 con el nordeste del grupo 7, más 1 municipio del grupo 6. Parte del grupo 4 se combina con parte del grupo 6. Dos municipios del grupo 4 se agregan a la combinación del grupo 2-8. Quedan así 5 subregiones como aparecen en el mapa de la figura 7.

Esta reclasificación es sometida nuevamente a un análisis discriminante para controlar la reclasificación, produciendo los resultados la matriz del cuadro 4 que muestra una clasificación perfecta.

FIGURA 7. Cinco subregiones de municipios en Pensilvania formados por análisis discriminante



CUADRO 1. Matriz para la evaluación de asignación de municipios a 5 subregiones

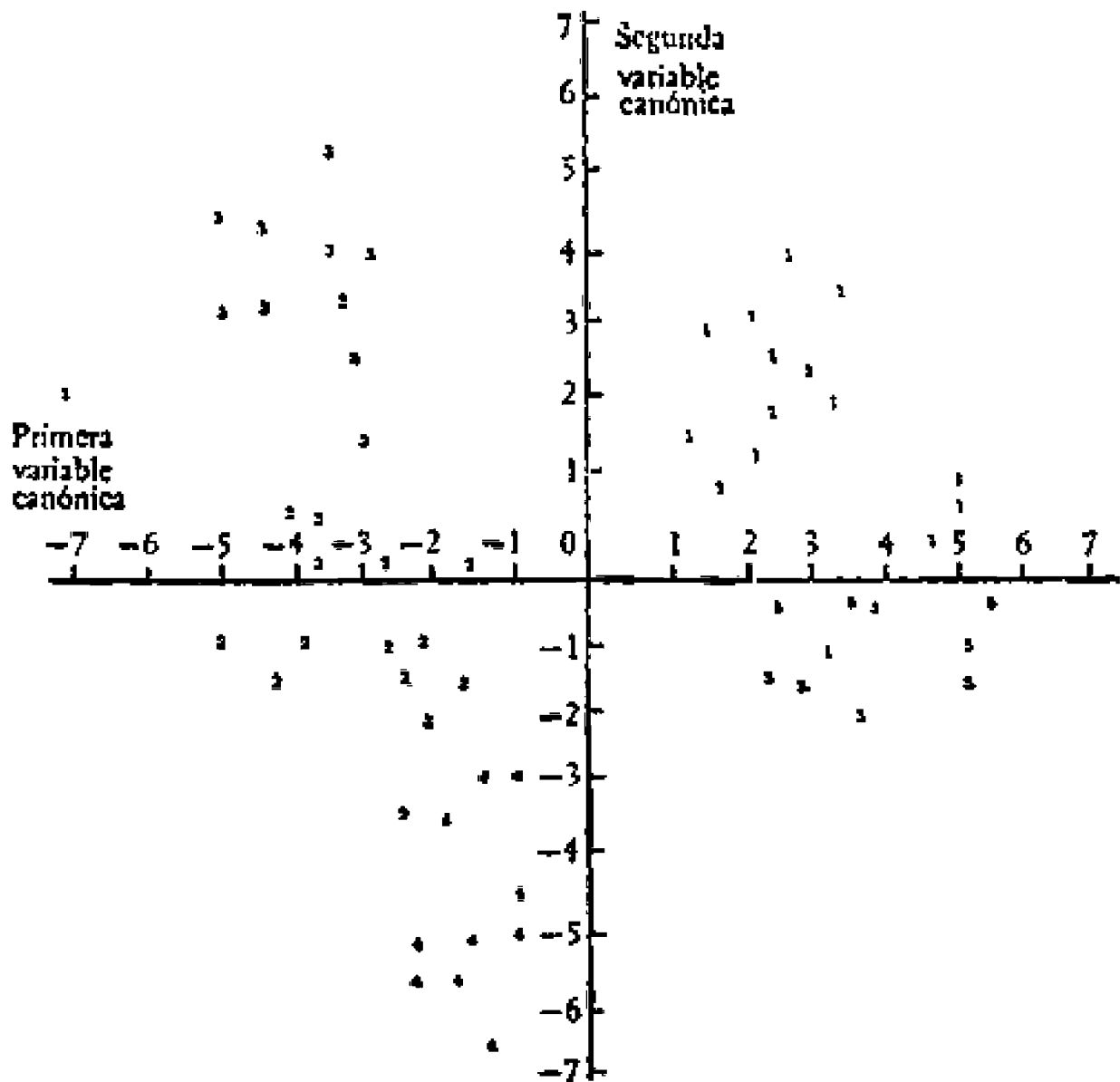
Subregión	Número de municipios por subregiones				
	1	2	3	4	5
1	16	0	0	0	0
2	0	14	0	0	0
3	0	0	11	0	0
4	0	0	0	11	0
5	0	0	0	0	13

Al mismo tiempo, la distribución gráfica de las nuevas variables canónicas muestra una completa separación entre los grupos, excepto para los grupos 1 y 5. Las variables canónicas de la figura 8 separan estos grupos.

Aunque resultante de una combinación de análisis discriminantes con técnicas de *cluster analysis*, se ve claramente cómo la técnica puede ser útil para aplicar criterios objetivos en la determinación de subregiones.

La técnica del análisis discriminante encuentra numerosas aplicaciones en el campo de la medicina (para la determinación de síndromes y eventualmente para diagnósticos más objetivos), en psicología, economía, etcétera.

FIGURA 8. Municipios agrupados en 5 subregiones



ANÁLISIS DE ESCALOGRAMA GUTTMAN (SUBPROGRAMA GUTTMAN SCALE)

La escala Guttman, conocida como “método del escalograma” o “análisis de escalograma”, tiene como objetivo definir lo más claramente posible qué es lo que está midiendo la escala, entendido esto como un problema de unidimensionalidad. Por el tipo especial de tratamiento al que se somete la escala, se busca la eliminación de factores extraños a las características o dimensión que se pretende medir.

La segunda característica de una escala Guttman es su propiedad de escala acumulativa, es decir que la respuesta positiva a un ítem supone que los ítems anteriores han sido respondidos en forma afirmativa. Se busca pues una coherencia en las pautas de respuesta de los sujetos, y esa coherencia es garantizada por medio de un coeficiente de

reproducibilidad. El tamaño del coeficiente (su valor máximo es 1.00) indica el grado en que la escala es acumulativa. En una escala cuya reproducibilidad es perfecta (1.00), las respuestas de los sujetos a todos los ítems pueden ser reproducidas con el solo conocimiento de la posición de rango. El escalograma Guttman combina aspectos de construcción utilizados por las escalas Lickert y Thurstone.⁷ A partir de una serie de ítems que son administrados a una muestra de sujetos que van a actuar como jueces, se procede a un análisis de los ítems en su conjunto, buscando la producción de una serie acumulativa de ítems. Los jueces se ordenan en término de los puntajes obtenidos en la escala, así como los ítems por grado de dificultad. Cuando la serie acumulativa de ítems es perfecta, todas las celdas cruzadas en el escalograma estarán en una posición sobre la diagonal que corre desde el ángulo superior izquierdo hasta el ángulo inferior derecho de la matriz. A más alto número de desviaciones de esta diagonal, mayor número de errores, es decir, menor reproducibilidad.

La idea, pues, es producir cortes y seleccionar ítems de manera tal que la reproducibilidad se maximice.

El cuadro que sigue muestra un análisis de escalograma con las respuestas de 20 jueces a 6 ítems, en términos de acuerdo-desacuerdo. Los acuerdos aparecen señalados con una x y los desacuerdos con un 0. Los valores encerrados en círculos son errores. (Cuadro 5.)

Las técnicas usadas para la determinación de los *cutting points*, esto es, los puntos que determinan el ordenamiento de los ítems, son básicamente dos (técnica de Cornell y técnica de Goodenough). La técnica de Cornell se realiza estableciendo puntos de separación en el orden de rango de los jueces, tales como se definirían éstos si la escala fuese perfecta. La técnica de Goodenough se basa en el cálculo de errores en base a pautas marginales.

Las posibilidades de establecimiento de *cutting-points* en escalas más complejas, esto es, con ítems con mayores categorías de respuesta que los presentados en el ejemplo anterior, obligan la mayoría de las veces a la reclasificación de categorías, simplificándolas. El cuadro 5 muestra un ejemplo algo más complejo.

CUADRO 5. Análisis de escalograma. Respuestas de 20 jueces a 6 ítems en términos de acuerdo-desacuerdo

Rango del juez	4	2	1	6	5	3	Puntaje
2	X	X	X	X	X	X	6
3	X	X	X	X	X	X	6
1	X	X	X	X	X	X	6
4	0	X	X	X	X	X	5
5	0	X	⓪	X	X	X	4
6	0	0	X	X	⓪	X	4
7	0	0	X	X	⓪	X	3
8	0	0	X	X	X	⓪	3
9	0	0	0	X	X	X	3
10	0	0	0	X	X	X	3
11	0	0	0	X	X	X	3
12	0	0	ⓧ	0	X	X	3
13	0	ⓧ	0	0	X	X	3
14	0	0	0	0	X	X	2
15	0	0	0	0	0	X	1
16	0	0	0	0	0	X	1
17	0	0	0	0	ⓧ	⓪	1
18	0	0	0	0	0	X	1
19	0	0	0	0	0	0	0
20	0	0	0	0	0	0	0
	3	6	8	11	14	16	

Como puede observarse, existen bastantes sujetos “fuera de posición” (en el ítem 1, hay varios sujetos que respondieron 4, y están por debajo de sujetos que respondieron 3; lo mismo hay sujetos que respondieron 1 en el ítem y tienen puntajes totales por encima de sujetos que dieron respuesta 2, 3 y 4). Los únicos sujetos que no tienen errores en este ítem son los que dieron respuesta 2. Se trata entonces de recombinar las categorías de respuesta de manera de minimizar los errores. Supongamos que las reclasificaciones propuestas para cada ítem, con sus respectivos pesos, son ahora las siguientes:

Ítem	Combinaciones	Nuevos pesos
1	4, 3 2, 1, 0	2 1 0
2	4, 3 2, 1, 0	2 1 0
3	4, 3 2, 1, 0	2 1 0
4	4, 3, 2 1, 0	2 1 0
5	4, 3 2, 1, 0	2 1 0
6	4, 3, 2 1, 0	2 1 0

Naturalmente, los puntajes totales de algunos sujetos variarán, alternando consecuentemente su orden de rango. El lector puede reconstituir la tabla.

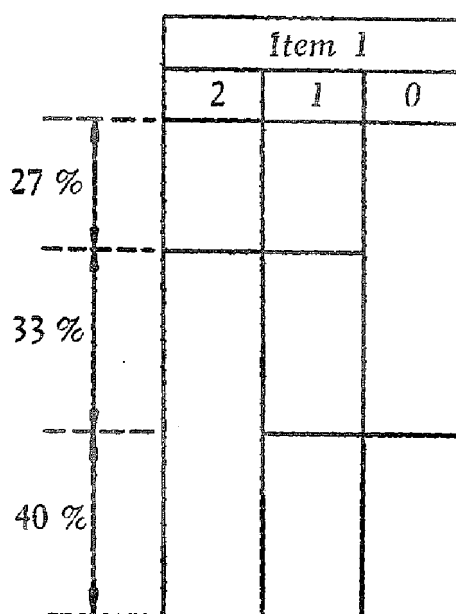
CUADRO 6. *Análisis de escalograma. Ordenación de puntajes y valores de respuesta en 6 ítems tipo escala Lickert*

Puntaje	Ítem 1					Ítem 2					Ítem 3					Ítem 4					Ítem 5					Ítem 6				
	4	3	2	1	0	4	3	2	1	0	4	3	2	1	0	4	3	2	1	0	4	3	2	1	0	4	3	2	1	0
24	X					X					X					X					X					X				
23	X					X					X					X					X					X				
21	X					X					X						X				X					X				
20		X				X					X						X				X					X				
19	X					X					X					X					X						X			
19	X					X					X					X					X					X				
19		X				X					X					X					X						X			
18	X					X					X					X					X					X				
18		X				X					X					X					X					X				
18		X				X					X					X					X					X				
18	X					X					X					X					X					X				
17		X				X						X					X				X					X				
16		X				X					X						X				X					X				
15		X						X	X								X				X					X				
15			X			X					X							X			X					X				
13		X						X			X							X								X	X			
12		X						X				X						X			X					X				
12		X						X				X					X	X			X					X				
11	X							X				X					X	X			X	X					X			
10			X					X				X					X				X					X				
8				X				X				X					X									X	X			
8				X				X				X					X				X					X				
8				X				X				X					X				X	X				X				
8			X					X				X					X				X					X				
7			X					X				X					X				X					X				
6				X				X				X					X				X					X				X
4			X					X				X					X				X					X				X
4				X				X				X					X				X					X				X
3				X				X				X					X				X					X				X
2				X				X				X					X				X					X				X

Los *cutting-points* en la técnica de Goodenough se determinan según la distribución de las respuestas de los jueces a cada una de las alternativas en los distintos ítems. Los cálculos para el cuadro 6 nos darían los siguientes valores:

	Ítems												
	1			2		3		4		5		6	
Categorías	2	1	0	2	0	2	0	2	0	2	0	2	0
Frecuencias	8	10	12	14	16	15	15	14	16	16	14	24	6
Porcentajes	27	33	40	47	53	50	50	47	53	47	53	80	20

Los *cutting-points* para cada uno de los ítems deben seguir entonces para cada ítem, los porcentajes respectivos. En el caso del ítem 1, los cortes serían entonces:



Es decir que en ítem 1, el primer corte o *cutting-point* caería entre el último sujeto con puntaje total 10 y el primero con puntaje total 9. Y así sucesivamente para el resto de los ítems.

Mediante estos *cutting-points*, se van a determinar pautas de respuestas correspondientes a cada corte. La pauta de respuesta es la manera "correcta" en que deberían distribuirse para cada puntaje total de cada juez si la escala tuviera perfecta escalabilidad. Cada respuesta que no sigue la pauta de respuesta ideal se considera un "error". En nuestro ejemplo, un juez con puntaje 8 debe seguir una pauta: 0-0-2-2-2-2. Un sujeto con puntaje 6 debe consecuentemente tener una pauta de respuesta 0-0-0-2-2-2. Es decir que primero hay que ordenar los ítems en términos de escalabilidad, para luego determinar las pautas correspondientes a cada valor. El principio de la técnica es bastante simple. Para 4 ítems dicotomizados y con valores de 1 y 0 para alternativas de respuesta "de acuerdo" y "en desacuerdo", respectivamente, un puntaje total de 3 puede seguir cuatro tipos de pautas diferentes:

- a) 0 1 1 1
b) 1 0 1 1
c) 1 1 0 1
d) 1 1 1 0

De estas pautas, solamente la primera es correcta, para ítems que están ordenados en forma acumulativa. Cada una de las pautas *b*, *c* y *d* tiene respectivamente 2 errores (uno por tener un 1 donde debería haber un 0, y otro por tener un 0 donde debería tener un 1). Para hacer más claro el método, vamos a dar un ejemplo más simple, en el que presentamos 4 ítems dicotomizados, las pautas de respuesta y el cálculo de los errores.

Items								Puntaje total	Pauta de respuesta	Errores			
1		2		3		4							
1	0	1	0	1	0	1	0						
X		X		X		X		4	1	1	1	1	0
	X	X		X		X		3					0
X			X	X		X		3					2
X		X		X			X	3	0	1	1	1	2
	X	X		X		X		3					0
	X	X		X		X		3					0
	X		X	X			X	2					2
	X		X	X		X		2	0	0	1	1	0
	X		X	X		X		2					0
	X		X		X	X		1					0
	X		X		X	X		1	0	0	0	1	0
	X		X		X		X	1					0
	X		X		X			0					0
	X		X		X		X	0	0	0	0	0	0
	X		X		X		X	0					0
Total errores									6				

El coeficiente de reproducibilidad se calcula según la fórmula ya conocida de:

$$r_p = 1 - \frac{\text{cantidad de errores}}{\text{número total de respuestas}}$$

El ejemplo arriba citado: $r_p = 1 - \frac{6}{60} = .90$.

El coeficiente de reproducibilidad nos indica la proporción de respuestas a los ítems que pueden ser correctamente reproducidas.

Presentamos ahora dos ejemplos, a partir de los cuales presentaremos criterios adicionales al coeficiente de reproducibilidad para la determinación del universo de contenido en la escala de Guttman. (Cuadro 7.)

Cálculo universo de contenido para el cuadro 7

	<i>Ejemplo A</i>	<i>Ejemplo B</i>
Coefficiente de reproducibilidad:	.949	1.00
Rango marginal mínimo:	.72	.60
Alcance de distribución marginal:	.229	.40

REFERENCIAS:

<i>Ejemplo A</i>	<i>Ejemplo B</i>
<i>Item 1:</i> Alejamiento de aguas servidas a pozo ciego y cámara séptica.	<i>Item 1:</i> Pileta de lavado.
<i>Item 2:</i> Agua corriente.	<i>Item 2:</i> Inodoro.
<i>Item 3:</i> Eliminador de residuos.	<i>Item 3:</i> Ducha.
<i>Item 4:</i> Techo de zinc, o loza o teja con aislación.	<i>Item 4:</i> Revestimiento de cemento o superior en baño o letrina.
<i>Item 5:</i> Piso de cemento o mejor.	
<i>Item 6:</i> Luz eléctrica.	

La *x* indica presencia del ítem en ambos ejemplos.

El coeficiente de reproducibilidad en el ejemplo *B*, cuyo valor es 1.00, nos permite decir, sin tener en cuenta la distribución gráfica, que el sujeto 18 tiene en su casa únicamente revestimiento de cemento, o mejor en su baño; mientras que el sujeto 33, cuyo puntaje es 3, tiene revestimiento, inodoro y ducha, pero no pileta de lavado.

CUADRO 7. Escalograma Guttman para la medición de nivel de vida determinada en función de materiales empleados en la construcción de viviendas y de las instalaciones sanitarias. Muestra de población de Colombres en el Departamento de Santa Cruz en Tucumán-Argentina.
Respuestas de 33 jefes de familia

Ejemplo A. Vivienda							Ejemplo B. Instalaciones sanitarias				
Sujetos	Ítems						Sujetos	Ítems			
	1	2	3	4	5	6		1	2	3	4
1	x	x	x	x	x	x	3	x	x	x	x
33	x	x	x	x	x	x	15	x	x	x	x
31	x	x	x	x	x	x	17	x	x	x	x
22	x	x	x	x	x	x	7	x	x	x	x
19	x	x	x	x	x	x	11	x	x	x	x
16	x	x	x	x	x	x	30	x	x	x	x
4	x	x	x	x	x	x	2	x	x	x	x
6	x	x	x		x	x	1	x	x	x	x
3		x	x	x	x	x	4	x	x	x	x
15		x	x	x	x	x	5	x	x	x	x
17		x	x	x	x	x	8		x	x	x
14		x	x	x	x	x	6		x	x	x
7		x	x	x	x	x	33		x	x	x
2		x	x	x	x	x	31		x	x	x
18		x	x	x	x	x	22		x	x	x
5		x	x	x		x	19		x	x	x
8		x	x		x	x	16		x	x	x
11		x	x	x		x	14		x	x	x
30		x	x	x		x	9				x
29			x		x	x	18				x
9				x	x	x	13				
13				x	x		25				
10				x	x		24				
26					x	x	29				
28					x	x	27				
24					x		26				
27						x	32				
20						x	12				
21						x	20				
25							28				
12							23				
23							10				
32							21				

Es muy difícil lograr escalabilidad perfecta, y consecuentemente, existen errores que

van a ser interpretados como errores de reproducibilidad. Guttman aconseja que los coeficientes de reproducibilidad no sean menos de .90.

El coeficiente de reproducibilidad (r_p) es un criterio necesario, pero no suficiente para la determinación de la escalabilidad de los ítems. Deben tomarse en cuenta otros factores. Stouffer *et. al.*,⁸ señalan cuatro criterios adicionales: a) Alcance de la distribución marginal; b) Pauta de errores; c) Número de ítems en la escala; y d) Número de categorías de respuestas.

a) Alcance de la distribución marginal

Es el más importante de los criterios adicionales, y debe acompañar al coeficiente de reproducibilidad. El criterio de distribución marginal es determinado por el rango marginal mínimo (MMR) que consiste en el r_p menos el promedio de los modos de las frecuencias relativas de las distribuciones de los ítems ($r_p - \text{MMR}$).

Para algunos, los valores de este criterio adicional deben variar entre .15 y .35; para otros el mínimo debe ser mayor que .10. Estos valores indican la *escalabilidad de los ítems*, dato que no es proporcionado por el r_p de manera completa; es decir, es posible alcanzar valores altos de r_p —digamos .90— y resultar una escalabilidad inaceptable. Éste es el caso en el cual los *cutting-points* están muy próximos entre sí, con el resultado de discriminar solamente en los extremos de la escala y no a lo largo de la misma. En nuestro ejemplo, los valores de r_p son altos y muy aceptables; los alcances de la distribución marginal, en cambio, son aceptables para el ejemplo A, y demasiado altos para el ejemplo B.

b) Pauta de errores

Cuando el r_p es menor que .90, pero es escalable, es decir que tiene un r_p MMR mayor que .10, estamos en presencia de más de una variable, mejor dicho, de una variable dominante y de otra u otras menores; en el área a través de la cual se ordenan los sujetos, este tipo de escalograma es denominado cuasi-escala. Éste no es el caso de los dos ejemplos que presentamos.

c) Número de ítems en la escala

A mayor número de ítems, mayor seguridad de que el universo, del cual estos ítems son una muestra, es escalable.

Por eso, cuando los ítems están dicotomizados, como es el caso en nuestros ejemplos, es aconsejable que su número sea mayor que 10. Pero puede usarse un número menor de ítems si las frecuencias marginales se colocan en un rango con recorridos del 30% al 70%.

En los ejemplos dados por nosotros, el rango de frecuencias es:

<i>Ejemplo A</i>		<i>Ejemplo B</i>	
Ítem 1	24%	Ítem 1	30%
Ítem 2	57%	Ítem 2	35%
Ítem 3	60%	Ítem 3	55%
Ítem 4	69%	Ítem 4	60%
Ítem 5	78%		
Ítem 6	87%		

Tenemos alguna seguridad, de acuerdo al requisito citado más arriba, de que el universo se comporta como la muestra.

d) Número de categorías de respuestas

Es otro criterio para asegurar la escalabilidad; cuanto mayor el número de categorías, mayor la seguridad de que el universo es escalable. Por ello, a pesar de la necesidad de reducir las categorías por razones prácticas (disminución del número de errores), hay que asegurarse de que tal reducción no es la resultante de obtener frecuencias marginales extremas (.90-.10) que, como vimos más arriba, no permiten errores, pero artificialmente.

Si mantenemos el número de alternativas de respuestas, a pesar de que aumentará el número de errores, disminuimos la posibilidad de que aparezca una pauta escalable cuando de hecho el universo no lo es.

El out-put del subprograma Guttman scale tiene la siguiente forma (suministramos los datos correspondientes al texto en el SPSS op. cit.):

Se trata de 3 ítems: 1) Miembro de organizaciones orientadas hacia el servicio. 2) Miembros de organizaciones ocupacionales. 3) Miembro de grupos recreacionales.

	<i>Recta. (3)</i>		<i>Org. Oc. (2)</i>		<i>M. O. S. (1)</i>		
	0	1	0	1	0	1	
	<i>Err</i> _____		<i>Err</i> _____		<i>Err</i> _____		
3	0	13	0	13	0	13	13
2	32	17	10	39	7	42	49
1	108	12	66	54	66	54	120
0	168	0	168	0	168	0	168
Sumas	360	42	244	106	241	109	350
%	88	12	70	30	69	31	
Errores	0	29	10	54	73	0	165

Coeficiente de reproducibilidad : .8419
 Reproducibilidad marginal mínima : .7552
 Porcentaje de mejoramiento : .0867
 Coeficiente de escalabilidad : .3541

Coeficientes de correlación

	Var. 1	Var. 2	Var. 3
Var. 1	1.000	.3496	.5151
Var. 2	.3496	1.000	.4024
Var. 3	.5151	.4024	1.0000
Escala			
Item	.2953	.2565	.8490

Los *cutting-points* se realizaron con una técnica similar a la Goodenough algo más simple que la especificada más arriba.

El cuadro se puede leer de la misma forma que el primero que presentamos más arriba. La diferencia es que aquí se representan en una doble columna los valores afirmativos (la x se remplaza con un 1); y los negativos (se representan con un 0). En la primera columna de la tabla figuran los valores 3, 2, 1, 0 que representan los posibles puntajes y que nos sirven para la determinación de los *cuttings-points* y por lo tanto para contar los errores. Un puntaje tres representa una pauta de respuesta 1 1 1; un puntaje dos representa una pauta 0 1 1 (siendo error por ejemplo 101 y 110); una pauta de respuesta de 1 es del tipo 0 0 1 (y no 100, 010). Examinando ahora el ítem que quedó ordenado como en primer lugar y que es el que más discrimina (solamente 12% de los sujetos dieron respuesta afirmativa), vemos que para el valor 0 del ítem no existe error pero sí hay 29 errores por debajo del *cutting point* (sujetos que han contestado afirmativamente el ítem y cuyo puntaje era menor que 3). En el ítem 2 existen 64 errores, 10 errores en 0 y 54 errores en 1; y en el ítem 3 hay 73 errores. El total de sujetos fue de 350, y es el resultado de la suma de la última columna, cuyos parciales indican la cantidad de sujetos que obtuvieron puntaje 3, 2, 1 y 0, respectivamente.

El coeficiente de reproducibilidad en el ejemplo es bastante bajo, y de haberse incluido más ítems en la escala, podría haberse eliminado algunos con el fin de aumentar la reproducibilidad, ya que la idea general es precisamente utilizar el escalograma para seleccionar ítems escalables.

El rango marginal mínimo nos da un valor que, restado del coeficiente de reproducibilidad, proporciona el porcentaje de mejoramiento o, como preferimos llamarle, alcance de la distribución marginal. En este caso el valor es de .1731, muy aceptable.

El coeficiente de escalabilidad corresponde a nuestra pauta de errores y en el ejemplo que presentamos representa un valor de .3541, muy bajo, e indica que la escala no es

unidimensional (los coeficientes deben ser mayores que .70).

Finalmente, las correlaciones que aparecen al fondo de la tabla son coeficientes Q de Yule para la correlación *inter* ítems, y coeficientes biserials para la intercorrelación de cada uno de los ítems con la suma del resto de los otros ítems. Con esto podemos analizar los ítems que no se correlacionan, ya sea con los otros ítems, ya sea con los valores de escala.

BIBLIOGRAFÍA

- Blalock, H.: *Social Statistics*; McGraw-Hill, Kogakusha, Tokio, 1972 (2ª ed.).
- Cacoullos, T. (ed.): *Discriminant Analysis and Applications*; Academic Press, Nueva York y Londres, 1973.
- Eisenbeis, R., y Avery, R.: *Discriminant Analysis and Classification Procedures*; Lexington Books, Mass., 1972.
- Fruchter, B.: *Introduction to Factor Analysis*; P. van Nostrand, Princeton, Nueva Jersey, 1954.
- Gansner, D., Seegrist, D., y Walton, G.: "Technique for Defining Subareas for Regional Analysis", en *Growth and Change*, octubre, 1971.
- Guilford, J. P.: *Psychometric methods*; McGraw-Hill, Nueva York, 1954.
- Gould, P. R.: "On the Geographical Interpretation of Eigenvalue", en *Institute of British Geographers*, núm. 42, diciembre, 1967 (pp. 53-85).
- Keerlinger, F.: *Foundations of Behavioral Research*; Holt, Rinehart and Winston, Nueva York, 1973 (2ª ed.).
- Nie, N., Hull C. H., Jenkins, J., Steinbrenner, K., y Bent D.: *Statistical Package for the Social Sciences*; McGraw-Hill, Nueva York, 1975 (2ª ed.).
- Padua, J.: *Escalas para la medición de actitudes*; CES, El Colegio de México, México, 1974.

X. LA PRESENTACIÓN DEL INFORME DE INVESTIGACIÓN

INGVAR AHMAN

EN ESTE último capítulo, nuestro interés se centra en la presentación del informe de investigación. Más exactamente, se trata de proponer diferentes pasos en una investigación, el modo en que pueden ser presentados y cuáles instrumentos metodológicos deberían ser incluidos. Creemos que la forma de una *checklist* es la más apropiada para estos fines. Sin embargo, si el estudio se hubiera efectuado de manera distinta de la forma expuesta en este libro, algunos de los puntos de la *checklist*, naturalmente, deben dejarse de lado y no considerarse; si hubo algún estudio tipo observación participante, por ejemplo, el cuestionario se cambia por una cédula estandarizada de registro de las observaciones; si el estudio fue de tipo más cualitativo y no existió la cédula, no se considera el punto, etcétera.

La idea, con respecto a esta lista, y tal como hemos venido insistiendo, es acentuar los siguientes puntos: *a)* la presentación, por parte del investigador, tanto de sus sistemas de hipótesis como de sus hallazgos, en una forma sistematizada; *b)* la inclusión de una descripción metodológica profunda que explique cómo se han obtenido los resultados; *c)* la presentación de los instrumentos utilizados en su forma original.

Con lo anterior, obtenemos los siguientes beneficios fundamentales: *a)* facilitar una simple inspección de los hallazgos obtenidos; *b)* facilitar la comparación con otros estudios efectuados en el mismo país y entre estudios en diferentes países; *c)* facilitar la replicación de la investigación completa, de baterías de preguntas y de escalas, así como la reformulación y mejoramiento de la calidad de posibles estudios en el futuro; *d)* facilitar, por las razones anteriores, la acumulación e incremento de hipótesis que ayuden en la formación y remodelación de la teoría.

La lista de control (check list)

1) El problema general. En este punto, el autor debe introducir la relación del problema tratado por la investigación con teorías específicas y generales. Se presenta, además, el sistema de hipótesis que se quiere investigar empíricamente. Esta presentación debe incluir abstracciones del sistema de hipótesis. Es decir, si es posible, resumirlo en forma de diagrama y fórmulas con símbolos. Asimismo, el análisis de otras investigaciones empíricas que han abordado el problema o aspectos de él, los resultados de estos estudios y su conexión con el modelo utilizado en la investigación que se comunica. Si el estudio actual fue influido por otro o si se han utilizado conceptos, modelos o instrumentos de otros estudios, éstos deben ser presentados aquí.

2) Diseño del estudio. El autor presentará el diseño que ha seleccionado para

investigar el problema, qué otras alternativas existían y por qué ha concluido por elegir la forma seleccionada. A esto hay que agregar la presentación de la técnica utilizada para la obtención de los datos y una breve explicación acerca de la unidad elegida. (Ejemplo: individuos masculinos de más de 20 años; familias; fábricas de tejidos; provincias; naciones; etcétera.)

En seguida, es necesario explicar las características de la población, incluyendo una descripción de las que han desempeñado algún papel de importancia en la selección de la muestra. Si el investigador, por razones económicas o de tiempo, ha extraído una muestra de la población, será necesario explicar el diseño de la muestra y las justificaciones que se han considerado para su selección y forma. El tamaño seleccionado y los problemas y limitaciones serán también de importancia.

3) *La descripción del grupo o de los grupos estudiados.* El autor debe presentar aquí la descripción de los grupos (o del grupo) estudiados, en base a conocimientos anteriores o pautas obtenidas en la investigación. Se trata de obtener una idea de las distribuciones de las variables de fondo para poder estimar mejor la parte explicativa del estudio. Las variables pueden ser de tipo ecológico, personal, contextual, etc. Además, una presentación descriptiva tiene, a veces, su propio valor, particularmente si el estudio está hecho en un campo desconocido o poco conocido.

4) *La medición de las variables investigadas.* Se concentra en las variables incluidas en el sistema de hipótesis y en la forma en que fueron medidas estas variables a través del instrumento utilizado en la colección de datos. Conviene empezar esta presentación con: a) las variables independientes: ¿cuáles fueron consideradas así en el esquema de análisis y qué forma de medición se aplicó? (Tipos de escalas que representan las variables: ¿nominales, ordinales, intervalos?) b) Las variables de “background”. ¿Qué se hizo en el diseño del estudio para controlar estas variables? ¿Cuáles fueron controladas?, etc. c) La presentación de la variable o las variables dependientes. ¿A qué nivel de medición (nominal, ordinal o intervalar) fue posible conseguir información en la variable? ¿Cuáles fueron las definiciones conceptuales? d) Estimación del grado de confiabilidad y validez en la medición de las variables en a, b y c. Si el investigador ha utilizado algunas técnicas o pruebas para comprobar el grado de confiabilidad y validez, es necesario incluir el tipo de prueba y una descripción del proceso. ¿Habrá otros métodos para verificar la confiabilidad y la validez en este tipo de estudios?

5) *Parte metodológica.* Se presenta una descripción del proceso utilizado para obtener los datos y de las facilidades y dificultades surgidas en el trabajo. ¿Cuáles fueron las instrucciones especiales dadas a los entrevistadores? ¿En qué manera se ha obtenido un incremento del grado de motivación con los entrevistados, frente a la entrevista misma? ¿Cuáles fueron, en general, las relaciones entre los entrevistadores y los entrevistados? ¿Cuáles fueron las relaciones entre la investigación y el público en general, antes, durante y después del estudio?

Debe presentarse, también, un análisis de la pérdida de información que ha sufrido el estudio. Un tipo de pérdida se refiere al grado de obtención de las entrevistas. Si la muestra fue extraída al azar, no se pueden sustituir personas que no pudieron ser

entrevistadas, por otras. Este último razonamiento es aplicable en aquellos casos en que los cuestionarios son enviados por correo u otro tipo de investigaciones especiales donde el porcentaje de rechazos en sí mismos refleja características que son significativas para el estudio. Debe anotarse el porcentaje de pérdida, así como algunos comentarios que justifiquen este porcentaje. Sin embargo, es posible construir diseños muestrales que contemplen casos de reservas que remplacen los rechazos. Otro tipo de pérdida puede referirse al grado de información perdida dentro de cada unidad (de cada entrevista). Ésta varía entre diferentes preguntas y diferentes baterías de preguntas y puede ser presentada y criticada por el autor. En esta parte es necesario incluir las relaciones del costo y tiempo necesario para efectuar el estudio.

Una parte puede presentar tipos de procesamiento de datos utilizados y una explicación de los tipos de análisis elaborados.

6) *Las relaciones encontradas.* La presentación de los hallazgos y relaciones encontradas en el estudio y sus relaciones con teorías conocidas.

Si se presenta el grado de relación entre las variables a través de *tests* estadísticos, debe también explicarse el tipo de *tests* aplicados y el grado de significancia existente. Se deben comparar los resultados obtenidos con los encontrados en otros estudios similares y, a la vez, ofrecer sugerencias al avance de una teoría o ideas para futuras investigaciones en el área.

7) *Apéndice.* Deben presentarse: a) Cuadros completos de correlaciones y otros tipos de asociaciones entre variables. Si fueron aplicados *tests*, se deben incluir los cuadros de los cruces y los resultados, con grados de significancia y tipos de *tests* aplicados. b) Una copia del cuestionario aplicado, acompañado de todos los “accesorios” utilizados, tales como las instrucciones dadas a los entrevistadores, hojas sueltas utilizadas en la entrevista, fotos y *tests*, etc. c) Una copia del código, especialmente si hay muchas preguntas abiertas en el cuestionario. Si el código se parece mucho a las alternativas que figuran en el cuestionario, no sería necesario, en tal caso, presentarlo. d) Mapas, diagramas completos y cualquier otro material de información. e) Un índice de autores y otro de materiales.

BIBLIOGRAFÍA GENERAL

- Adorno, Frenkel-Brunswick, *et al.*: *The Authoritarian Personality*; Harper & Brothers, Nueva York, 1950.
- Bachelard, G.: *La formación del espíritu científico. Contribución a un psicoanálisis del conocimiento objetivo*; Siglo XXI, Buenos Aires, 1972.
- Barton, A.: "The Concept of Property-Space in Social Research", en Lazarsfeld, P., y Rosenberg, M.: *The Language of Social Research*, The Free Press of Glencoe, Illinois, 1955.
- Beveridge, W.: *The Art of Scientific Investigation*; W. W. Borton Co., Nueva York, 1957.
- Blalock, H.: *Social Statistics*; McGraw-Hill Kogakusha, Tokio, 1972, 2ª ed.
- : *Causal Inferences in Nonexperimental Research*; The University of North Carolina Press, Chapel Hill, 1964.
- : *Theory Construction*; Prentice-Hall, Nueva Jersey, 1969.
- (ed.): *Causal Models in the Social Sciences*; Aldine-Altherton, Chicago, 1971.
- Blalock, H., y Blalock, A.: *Methodology in Social Research*; McGraw-Hill, Nueva York, 1968.
- Boudon, R., y Lazarsfeld, P.: *Metodología de las Ciencias Sociales*; Laia, Barcelona, 1973.
- : *El análisis empírico de la causalidad*, Laia, Barcelona, 1974.
- Bourdieu, P., Passeron, J-C., Chamboredon, J-C.: *Le Métier de Sociologue*; Mouton Bordas, París, 1968.
- Borko, H. (ed.): *Computer Applications in the Behavioral Sciences*; Prentice-Hall, Nueva Jersey, 1962.
- Brodbeck, M.: *Readings in the Philosophy of the Social Sciences*; The Macmillan Co., Nueva York, 1968.
- Brown, R.: *Psicología social*; Siglo XXI, México, 1972.
- Bunge, M.: *Causalidad: el principio de la causalidad en la ciencia moderna*; Eudeba, Buenos Aires, 1965.
- : *La ciencia, su método y su filosofía*; Siglo XX, Buenos Aires, 1970.
- Buss, A.: *Psychopathology*; J. Wiley and Sons, Nueva York, 1966.
- Bustamante, J.: *Mexican Immigration and the Social Relations of Capitalism*. Doctoral Dissertation; University of Notre-Dame, N.-D., Indiana, 1974.
- Cacoulios, T. (ed.): *Discriminant Analysis and Applications*; Academic Press, Nueva York y Londres, 1973.
- Castell, M.: *Problemas epistemológicos implícitos en la práctica sociológica* (mimeo), FLACSO, Santiago de Chile, 1968.

- Cochran, W.: *Sampling Techniques*; J. Wiley and Sons, Nueva York, 1953.
- Coleman, J.: *Introduction to Mathematical Sociology*; The Free Press of Glencoe, Nueva York, 1964.
- Coombs, C. H.: *A Theory of Psychological Scaling*; Engineering Research Institute, University of Michigan, Ann Arbor, Mich., 1952.
- Chevry, G.- R.: *Pratique des enquêtes statistiques*; Presses Universitaires de France, 1962.
- Christie, R. y Jahoda, M. (ed.): *Studies in the Scope and Methods of "The Authoritarian Personality"*; The Free Press, Nueva York, 1954.
- Duverger, M.: *Métodos de las ciencias sociales*; Ariel, Barcelona. Traducción del original francés, *Méthodes des Sciences Sociales*, Presses Universitaires de France, 1961.
- Denzin, N.: *Sociological Methods. A sourcebook*; Aldine, Chicago, 1970.
- Edwards, A.: *Techniques of Attitude Scale Construction*; Appleton-Century-Crofts, Nueva York, 1957.
- Eisenbeis, R., y Avery, R.: *Discriminant Analysis and Classification Procedures*; Lexington Books, Mass., 1972.
- Festinger, L., y Katz, D.: *Research Methods in the Behavioral Sciences*; The Dryden Press.
- Forcese, D. y Richer, S. (ed.): *Stages of Social Research*.
- Fruchter, B.: *Introduction to Factor Analysis*; P. van Nostrand, Princeton, Nueva Jersey, 1954.
- Galtung, J.: *Teoría y métodos de la investigación social* (dos tomos); Eudeba, Buenos Aires, 1966.
- Gansner, D., Seegrist, D., y Walton, G.: "Technique for Defining Subareas for Regional Analysis", *Growth and Change*, octubre de 1971.
- Germani, G.: *Política y sociedad en una época de transición*; Paidós, Buenos Aires, 1968.
- Gerth, H., y Mills, N. C.: *Carácter y estructura social*; Paidós, Buenos Aires, 1963.
- Goode, W. J., y Hatt, P. K.: *Métodos en Pesquisa Social*, Companhia Editora Nacional, São Paulo, 1960. Traducción del original inglés "Methods on Social Research", McGraw-Hill Co. Inc., Nueva York, 1972.
- Gould, P. R.: "On the Geographical Interpretation of Eigehvalue", en *Institute of British Geographers*, num. 42, diciembre de 1967 (págs. 53-85).
- Greenwood, E.: *Sociología experimental*, Fondo de Cultura Económica, México, 1951. Traducción del original inglés, *Experimental Sociology: A study in Method*; King's Grown Press, 1945.
- Guilford, J. P.: *Psychometric Methods*; McGraw-Hill, Nueva York, 1954.
- Hays, W.: *Statistics for Psychologists*; Holt, Rinehart and Winston, Nueva York, 1963.
- Heintz, P.: *Curso de sociología*; Eudeba, Buenos Aires, 1965.
- Hempel, C.: *Fundamentals of Concept Formation in Empirical Science*; International Encyclopedia of Unified Science, The University of Chicago Press, vol. II, núm. 7,

- Chicago, 1969 (10ª impr.).
- Hovland, C., *et al.*: "A Baseline for the Measurement of Percentage Change", en Lazarsfeld, P., y Rosenberg M.: *The Language of Social Research*; The Free Press of Glencoe, Ill., 1955.
- Hyman, H.: "Survey Design and Analysis". The Free Press, Publishers, Glencoe, Ill., 1957.
- : *Secondary Analysis of Sample Surveys. Principles, Procedures and Pottentialities*; J. Wiley and Sons, Nueva York, 1972.
- Hyman, H., *et al.*: *Interviewing in Social Research*; The University of Chicago Press, Chicago, 1965.
- Jahoda, M.: *Current Concepts of Positive Mental Health*; Basic Books, Nueva York, 1958.
- Kahn, R., y Cannell, Ch.: *The Dynamics of Interviewing: Theory, Techniques and Cases*; J. Wiley and Sons, Nueva York, 1962.
- Kaplan, A.: *The Conduct of Inquiry. Methodology for Behavioral Sciences*; Chandler, San Francisco, 1964.
- Kerlinger, F.: *Foundations of Behavioral Research*; Holt, Rinehart and Winston, Nueva York, 1973 (2ª ed.).
- Kretch, D., y Gruschfield, R. S.: *Theory and Problems of Social Psychology*; McGrawHill, Nueva York, 1948.
- Krimerman, L.: *The Nature & Scope of Social Science: A critical Anthology*; Appleton-Century-Crofts, Nueva York, 1969.
- Kuhn, T.: *The Structure of Scientific Revolutions*, International Encyclopedia of Unified Science, vol. II, núm. 2; The University of Chicago Press, Chicago, 1970 (2ª ed.).
- Lazarsfeld, P., y Rosenberg, M.: *The Language of Social Research*; The Free Press of Glencoe, Illinois, 1955.
- Lazarsfeld, P.: *Qualitative Analysis. Historical and Critical Essays*; Allyn and Bacon, Boston, 1972.
- : "Evidence and Inference in social research", en *Evidence and Inference*, Lerner *et al.*; The Free Press, Nueva York, 1958.
- Lerner *et al.*: *Evidence and Inference*; The Free Press, Nueva York, 1958.
- Lundberg, G. A.: *Técnica de la investigación social*, FCE, México, 1954.
- McCormick, T.: *Técnica de la estadística social*, FCE, México, 1954. Traducción del original inglés, *Elementary Social Statistics*, McGraw-Hill Book Co. Inc., Nueva York, 1941.
- McCall, G., y Simmons, J.: *Issues in Participant Observation. A text and a Reader*; Addison-Wesley, Massachusetts, 1969.
- Merton, R.; Broom, y Cottrell: *Sociology Today*; Basic Books, Nueva York, 1959.
- Merton, R.: *Teoría y estructuras sociales*; Fondo de Cultura Económica, México, 1964.
- Miller, D. C.: *Handbook of Research Design and Social Measurement*; D. McKay Co., Nueva York, 1964.
- Moser, C. A.: *Survey Methods in Social Investigation*; Heinemann, Londres, 1957.

- Nettler, G.: *Explanations*; McGraw-Hill, Nueva York, 1970.
- Newcomb, T. M.: *Personality and Social Change: Attitude Formation in a Student Community*; Dryden, Nueva York, 1943.
- Nie, N., Hull, C. H., y Jenkins, J.: *Statistical Package for the Social Sciences*; McGraw-Hill, Nueva York, 1975 (2ª ed.).
- Padua, J., Quevedo, S., Faria, R., y Ochoa, J.: *Estudio del programa de asistencia escolar proporcionado por la JNAEB*; ELAS-JNAEB, Santiago de Chile, 1967.
- Padua, J.: *Estudios globales y estudios parciales: un análisis desde la perspectiva epistemológica* (mimeo); CES, El Colegio de México, México, 1974.
- : *Some Problems in Logical and Sociological Empiricism* (mimeo); CES, El Colegio de México, México, 1973.
- Parsons, T., y Shieds, E.: *Toward a General Theory of Action*; Harper & Row, Nueva York, 1951.
- Parten, M.: *Survey, Polls, and Samples*, Harper & Row, Nueva York, 1950.
- Riley, M.: *Sociological Research: A Case Approach*; Harcourt, Brace and World, Nueva York, 1963.
- Robinson, J. y Shever, Ph.: *Measures of Social Psychological Attitude*; Survey Research Center, The University of Michigan, Ann Arbor, Michigan, 1971 (3ª imp.).
- Robinson, J.; Rusk, J., y Head, K.: *Measures of Political Attitudes*; Survey Research Center, ISR, The University of Michigan, Ann Arbor, 1969 (2ª imp.).
- Robinson, J., Athenasion, R., y Head, K.: *Measures of Occupational Attitudes and Occupational Characteristics*; Survey Research Center, ISR, The University of Michigan, Ann Arbor, Michigan, 1969.
- Ruiz, M., y Nieto, D.: *Tratado elemental de botánica*; Porrúa, México, 1950.
- Selltiz, C., Jahoda, M., Deutsch, M., y Cook, S.: *Research Methods in Social Relations*; Holt, Rinehart and Winston, Nueva York, 1959.
- Siegel, S.: *Nonparametric Statistics for the Behavioral Sciences*; McGraw-Hill, Nueva York, 1956.
- Simon, J.: *Basic Research Methods in Social Science: The Art of Empirical Investigation*; Random House, Nueva York, 1969.
- Stinchcombe, A.: *Constructing Social Theories*; Harcourt, Brace, and World, Nueva York, 1966.
- Stouffer, S., et al.: *Studies in Social Psychology in World War II. Measurement and Prediction* (vol. IV); J. Wiley y Sons, Nueva York, 1966.
- Thomas, W. y Znaniecki: *The Polish Peasant in Europe and America*.
- Thurstone, L.: *The Measurement of Attitudes*; The University of Chicago Press, Chicago, 1929.
- Torgerson, W.: *Theory and Methods of Scaling*; J. Wiley and Sons, Nueva York, 1958.
- Unión Panamericana. Naciones Unidas: *Manual sobre métodos para la elaboración de datos*. Parte I (Edición Provisional. Preparado por la Oficina de Estadística de las Naciones Unidas para la Agricultura y la Alimentación, Roma, Italia. Traducido y publicado por el Departamento de Estadística de la Unión Panamericana, 1959).

Weber, M.: *Economía y Sociedad* Fondo de Cultura Económica, México, 1969 (2ª ed.); (1ª reimp.).

Zeisel, H.: *Dígalos con números*; Fondo de Cultura Económica, México, 1974.

Zetterberg, H.: *On Theory and Verification in Sociology*; The Bedminster Press, Nueva York, 1965 (3ª ed.).

I. LA ORGANIZACIÓN DE UN “SURVEY”

¹ A este respecto, afortunadamente se está generalizando en América Latina la existencia de centros de documentación con personal técnico especializado en clasificación, ordenación, etc., de material de documentación por área, subáreas, etc., que facilitan enormemente la tediosa tarea para el investigador de seleccionar el material relevante.

² La documentación descriptiva debe ser sometida a crítica antes de aceptar las conclusiones a que en ella se pueda llegar. Sin caer en el hipercriticismo, los documentos deben ser examinados tanto en lo relativo a sus características internas como externas, realizando un análisis de las fuentes citadas por los documentos, de los intereses que puedan estar modulando la selección de fuentes, capacidad del observador para seleccionar los aspectos más relevantes de la situación, o solamente de aquellos que coincidan con su punto de vista. Dependiendo del tema de investigación, y sobre todo cuando la documentación es de tipo descriptivo, conviene realizar algunos análisis comparativos entre distintas fuentes con el fin de comprobar la existencia de contradicciones.

³ Bustamante, por ejemplo, para realizar su estudio sobre los emigrantes mexicanos a los Estados Unidos (ver Bustamante, J.: *Mexican Immigration and the Social Relations of Capitalism*. Doctoral Dissertation [sociología]; University of Notre-Dame, Notre-Dame, Indiana, 1974) entró a la frontera de los Estados Unidos como “espalda mojada”, esto es, como emigrante “ilegal”; fue detenido en Texas, encarcelado y enviado de vuelta a México, teniendo oportunidad de vivir el proceso en su totalidad. Entrevistó a braceros y adquirió así una gran experiencia en el terreno, lo cual dio a su trabajo una riqueza imposible de lograr por el simple análisis abstracto de la situación.

⁴ Hyman, H., *et al.*, *Interviewing in Social Research*; The University of Chicago Press, Chicago, 1965.

⁵ Existe buena bibliografía especializada sobre el tema; entre otros recomendamos: Kahn, R., y Cannell, Ch.: *The Dynamics of Interviewing: Theory, Techniques and Cases*; John Wiley & Sons, Nueva York, 1962.

II. EL PROCESO DE INVESTIGACIÓN

¹ Ésta es la diferenciación que hace Galtung en *Teoría y métodos de la investigación social*; Eudeba, Buenos Aires, 1966. Nosotros preferimos diferenciar los estudios en descriptivos y explicativos según se basen en definiciones o proposiciones.

² La realidad no se da en el análisis sino que es reconstruida en el proceso analítico mismo. No medimos *la* realidad, sino que medimos *en la* realidad, mediante la traducción del “objeto real” en un “objeto científico”, a través de un haz de conceptos que tienen significados nominales para una teoría general o particular (ver Castell, M.: *Problemas epistemológicos implícitos en la práctica sociológica-FLACSO*; Santiago de Chile, 1968).

³ Riley, M.: *Sociological Research: A Case Approach*; Harcourt, Brace & World, Nueva York, 1963.

⁴ El lector interesado en el problema puede recurrir a una abundante bibliografía especializada, entre otros recomendamos: Brodbeck M. (ed.): *Readings in the Philosophy of the Social Sciences*; The McMillan Co., Nueva York, 1968; Krimmerman, L.: *The Nature and Scope of Social Science: A Critical Anthology*; Appleton-Century-Crofts, Nueva York, 1969; Kaplan, A.: *The Conduct of Inquiry. Methodology of Behavioral Sciences*; Chandler Publ. Co., San Francisco, 1964; Bachelard, G.: *La formación del espíritu científico. Contribución a un psicoanálisis del conocimiento objetivo*; Siglo XXI, Buenos Aires, 1972; Bourdieu, P., Passeron, J. C., Chamboredon, J. C.: *Le Métier de Sociologue*; Mouton Bordas, París, 1968.

⁵ Zetterberg H.: *On Theory and Verification in Sociology*; The Bedminster Press Nueva York, 1965 (3ª ed.).

⁶ Piaget, J.: “La situación de las ciencias del hombre dentro del sistema de las ciencias”, en Piaget, J., Mackenzie, W., Lazarsfeld P., y otros, *Tendencias de la investigación en las ciencias sociales*; Alianza Universidad, UNESCO, 1970.

⁷ Bunge, M.: *La ciencia, su método y su filosofía*; Siglo XX, Buenos Aires, 1970; y *Causalidad: el principio de la causalidad en la ciencia moderna*; Eudeba, Buenos Aires, 1965.

⁸ Las preguntas que Zetterberg (*op. cit.*) sugiere son: 1) ¿existe un orden subyacente a la realidad social? 2) Si es así, ¿se han descubierto algunas leyes sociológicas? 3) Si es así, ¿se han combinado esas leyes en teorías que explican la realidad social? 4) Si es así, ¿se han utilizado esas teorías para calcular soluciones a problemas prácticos?

⁹ La naturaleza de la explicación en ciencias sociales involucra no solamente una serie de argumentos sobre el estatuto de la explicación como conocimiento científico (ver, por ejemplo, Padua, J.: “Estudios globales y estudios parciales: un análisis desde la perspectiva epistemológica”, ponencia presentada en el Seminario sobre Interrelaciones entre la Dinámica Demográfica y la Estructura y Desarrollo Agrícola, en Cuernavaca, México, noviembre de 1974), sino que tiene que ver con el tipo de explicación más apropiada a los fenómenos sociales (teleológica, cuasiteleológica, no teleológica, probabilística, genética-estructural, etc.). No entramos a considerar estos argumentos en este texto, pero recomendamos parte de la bibliografía que se cita en la primera parte de este capítulo, con el agregado tal vez del texto de Nagel, E.: *The Structure of Science. Problems in the Logic of Scientific Explanation*; Harcourt, Brace & World, Nueva York, 1961.

¹⁰ Para una discusión en detalle del tema conviene leer, de Hempel, C.: “*Fundamentals of Concept Formation in Empirical Science*”; International Encyclopedia of Unified Science., vol. II, núm. 7; y de Blalock, H.: *Theory Construction*; Prentice-Hall, Nueva Jersey, 1969.

¹¹ Beveridge, W.: *The Art of Scientific Investigation*; W. W. Norton Co., Nueva York, 1957.

¹² Para un tratamiento muy penetrante del tema hipótesis, recomendamos la lectura del tomo II de Galtung (*op.*

cit.) que contiene uno de los tratamientos más completos, del cual señalamos aquí solamente algunos puntos relevantes a nuestra línea de argumentación.

¹³ Un tratamiento bastante extenso del tema puede encontrarse en Zetterberg (*op. cit.*).

¹⁴ Para un buen desarrollo de lo que se conoce como “serendipiti” ver Merton, R.: *Teoría y estructuras sociales*; FCE, México, 1964.

¹⁵ Adorno, F.-B., *et al.*: *The Authoritarian Personality*; Harper and Brothers, Nueva York, 1950.

¹⁶ Lazarsfeld P.: “Evidence and Inference in Social Research”, en *Evidence and Inference* Lerner *et al.*; The Free Press, Nueva York, 1958.

¹⁷ Existe una abundante bibliografía crítica sobre los estudios de “La personalidad autoritaria”; ver, por ejemplo, un buen resumen en Brown, R.: *Psicología social*; Siglo XXI, México, 1972. Para un análisis más profundo ver, por ejemplo, Christie, R., & Jahoda, M. (eds.): *Studies in the Scope and Methods of “The Authoritarian Personality”*; The Free Press, Nueva York, 1954.

¹⁸ La palabra teoría es utilizada aquí en un sentido bastante laxo, para referirnos más que nada al sistema de racionalidad que los psicólogos utilizan para “explicar” el origen, naturaleza y desarrollo de la enfermedad (sería más exacto hablar de aproximaciones psicoanalíticas, del aprendizaje, biológicas, etc.). Lo mismo ocurre en sociología, donde las taxonomías están todavía en estado de aceptación parcial, según la escuela de preferencia de algunos autores (hablando de teorías del conflicto, funcionalistas, etcétera).

¹⁹ Buss, A.: *Psychopathology*; John Wiley and Sons, Nueva York, 1966.

²⁰ Jahoda, M.: *Current Concepts of Positive Mental Health*; Basic Books, Nueva York, 1958.

²¹ Buss (*op. cit.*) incluye además en la lista trastornos en los niños, inteligencia subnormal y daños en el cerebro.

²² Las taxonomías más completas pueden verse en Weber, M.: *Economía y sociedad* (2 tomos); FCE, México, 1969 (1ª reimpr.); y en Parsons, T., y Shieds, E. (ed.). *Toward a General Theory of Action*; Harper & Row, Nueva York, 1951.

²³ Germani, G.: *Política y sociedad en una época de transición*; Paidós, Buenos Aires, 1968.

III. MUESTREO

¹ Galtung, J.: *Teoría y métodos de la investigación social*; Eudeba, Buenos Aires, 1966.

² Lazewitz, B.: "Sampling Theory and Procedures"; en Blalock, H. y Blalock, A., (eds.): *Methodology in Social Research*; McGraw-Hill, Nueva York, 1968. Cochran, W.: *Sampling Techniques*; J. Wiley & Sons, Nueva York, 1953. Yates, F.: *Sampling Methods for Censuses and Surveys*; Griffin, Londres, 1953.

³ Algunos autores, como Galtung, no incluyen este tipo de muestra entre las probabilísticas, ya que los procedimientos de muestreo rompen el principio de aleatoriedad. Este mismo problema surge en el caso de las muestras sistemáticas, en la medida en que, una vez seleccionada la primera unidad, la probabilidad de las unidades siguientes es cero o uno, es decir que luego de seleccionada la primera unidad la muestra es finalista.

⁴ Al lector no familiarizado con los problemas de la estadística inferencial recomendamos nuestro capítulo sobre conceptos estadísticos básicos, o algún texto sobre estadística; recomendamos especialmente: *Social Statistics* de Blalock, H.; McGraw-Hill Kogakusha, Tokio, 1972. (Ed. en español: *Estadística social*; Fondo de Cultura Económica, México, 1979.)

⁵ Galtung, *op. cit.*, pp. 65 a 67.

⁶ Tabla extraída de Galtung, *op. cit.*, p. 66.

⁷ Extraída de Blalock, H.: *Social Statistics*; International Student Edition, McGraw-Hill Kogakusha, Tokio, 1972, p. 555.

⁸ Uno puede tomar otro criterio, por supuesto, para definir los conglomerados: región geográfica, cantidad de habitantes, etc. Hay que tratar que éstos sean lo más homogéneos posible. En los últimos *escalones* vamos a utilizar criterios diferentes de las divisiones políticas.

⁹ Galtung, J.: *op. cit.*, pp. 61 y ss.

IV. EL CUESTIONARIO

¹ Ver capítulos “La organización de un *survey*” y “El trabajo de campo”.

² Existen algunas publicaciones especializadas, tales como las del Institute of Social Research de la Universidad de Michigan, muy útiles para la formulación de preguntas o escalas en un cuestionario, además de revistas periódicas tales como: *American Sociological Review*; *American Journal of Sociology*; *Social Forces*; *Psychological Bulletin*, etcétera.

³ Una lista de control para recordar los objetivos que uno tiene que cumplir en una tarea determinada.

⁴ Ejemplo tomado de Duverger, M.: *Métodos de las ciencias sociales*; Ariel, Barcelona, 1961.

⁵ Para detalles de la muestra y del estudio, ver: Faria, R., Ochoa, J., Quevedo, S., y Padua, J.: *Estudio del Programa de Asistencia Alimenticia Proporcionado por la Junta Nacional de Auxilio Escolar y Becas*; ELAS-JNAEB, Santiago de Chile, 1967; y *El rendimiento escolar: un análisis en base a algunas variables estructurales*; ELAS-JNAEB, Santiago de Chile, 1968.

V. LA CODIFICACIÓN

* Quiero aprovechar esta oportunidad para señalar mi deuda intelectual con el profesor doctor don Jorge Garmendia, no sólo en lo referente a la confección de este trabajo —sin cuyas sugerencias y alientos no hubiera sido realizado—, sino en cuanto a mi formación profesional en la metodología de la investigación social.

¹ Duverger, M.: *Métodos de las ciencias sociales*; Ariel, Barcelona. Traducción del original francés, *Méthodes des Sciences Sociales*; P.U.F., 1961.

² Hyman, H.: *Survey Design and Analysis*; The Free Press, Publishers, Glencoe, Ill., 1957.

³ Selltiz, C., Jahoda, M., Deutsch, M., y Cook, S.: *Research Methods in Social Relations*; Holt, Rinehart and Winston, Nueva York, 1959.

⁴ También podemos señalar la definición de Chevy: “La codificación de los cuestionarios es la operación que consiste en traducir, utilizando los códigos establecidos para ese fin, las respuestas literales o numéricas en indicaciones cifradas que facilitarán la clasificación; las cuales podrán inmediatamente ser transcritas en perforaciones de una columna de una tarjeta perforada”. “Pero puede ser útil también cuando se encara una simple explotación manual.” Chevy, G.-R.: *Pratique des enquêtes statistiques*; Press Universitaires de France, 1962.

⁵ Para una visión general de los símbolos que pueden emplearse para codificar, véase el apéndice.

⁶ Chevy, G., *op. cit.*

⁷ Duverger, M., *op. cit.*

⁸ Duverger, M., *op. cit.*

⁹ Hyman, H., *op. cit.*

¹⁰ Duverger, *op. cit.*

¹¹ Selltiz, C., Jahoda, M., *et al.*, *op. cit.*

¹² Chevy, G., *op. cit.*

¹³ Hyman, H., *op. cit.*

¹⁴ Chevy, G., *op. cit.*

¹⁵ Selltiz, C., Jahoda, M., *et al.*, *op. cit.*

¹⁶ Chevy, G., *op. cit.*

¹⁷ Chevy, G., *op. cit.*

VI. ESCALAS PARA LA MEDICIÓN DE ACTITUDES

¹ Coombs, C. H.: *A Theory of Psychological Scaling*; Engineering Research Institute, University of Michigan, Ann Arbor, Mich., 1952.

² Un ejemplo mucho más claro sería preguntar en qué país nació, o cuál es su sexo.

³ Stouffer, S., *et al.*: *Studies in Social Psychology in World War II. Measurement and Prediction* (vol. IV); J. Wiley & Sons, Nueva York, 1966.

⁴ Kretch D., y Cruschfield, R. S.: *Theory and Problems of Social Psychology*; McGraw-Hill, Nueva York, 1948.

⁵ Newcomb, T. M.: *Personality and Social Change: Attitude Formation in a Student Community*; Dryden, Nueva York, 1943.

⁶ Thomas, W., y Znaniecki, F.: *The Polish Peasant in Europe and America*.

⁷ Edwards, A.: *Techniques of Attitude Scale Construction*; Appleton-Century-Crofts, Nueva York, 1967.

⁸ Para llegar al promedio sumamos la columna correspondiente al ítem 1 en los 13 sujetos con puntajes altos dividiéndolos por el total de casos: $(4+4+3+4+4+4+4+3+4+3+3+4+4)/13$. En forma idéntica, para el grupo bajo el cálculo es: $(2+0+2+1+0+1+0+1+2+3+0+0+0)/13$.

⁹ Stouffer, S., *et al.*: *Studies in Social Psychology in World War II. Measurement and Prediction*, vol. IV; J. Wiley & Sons, Nueva York, 1966.

VII. EL TRABAJO DE CAMPO

¹ *R* = Respondente.

VIII. ANÁLISIS DE DATOS: EL CONCEPTO DE PROPIEDAD-ESPACIO Y LA UTILIZACIÓN DE RAZONES, TASAS, PROPORCIONES Y PORCENTAJES

¹ Consideramos en este capítulo únicamente razones, tasas, proporciones y porcentajes. En el capítulo de Jorge Padua se consideran modelos estadísticos más sofisticados y para niveles de medición ordinales e intervalares, además de los nominales.

² NES = Nivel socioeconómico.

³ Para más detalles sobre el tema consultar: Barton, A.: “The Concept of Property-Space in Social Research”, en Lazarsfeld, P., y Rosenberg, M.: *The Language of Social Research*; The Free Press, Nueva York, 1955.

⁴ Gerth, H., y C. W. Mills.: *Carácter y estructura social*; Paidós, Buenos Aires, 1963.

⁵ Heintz, P.: *Curso de sociología*; Eudeba, Buenos Aires, 1965.

⁶ Barton, A.: *op. cit.*

⁷ Zeigel, H.: *Digalo con números*. Fondo de Cultura Económica, México, 1961.

⁸ Hovland, C. J., et. al.: “A Baseline for the Measurement of Percentage Change”, en Lazarsfeld, P., y Rosenberg, M.: *The Language of Social Research*; The Free Press, Nueva York, 1965.

IX. ANÁLISIS DE DATOS: PAQUETE ESTADÍSTICO PARA LAS CIENCIAS SOCIALES (SPSS): OFERTA Y CONDICIONES PARA SU UTILIZACIÓN E INTERPRETACIÓN DE RESULTADOS

¹ Guilford, J. P.: *Psychometric Methods*; McGraw-Hill, Nueva York, 1954.

² Blalock, H.: *Social Statistics*; McGraw-Hill, Kogakusha, Tokio, Japón, 1972 (2ª ed.).

³ Salvo en el caso de utilizar la correlación parcial para predicciones en la utilización de regresiones en cuyo caso se acostumbra a interpretar $r_{12.3}$, denotando 1 la variable dependiente, 2 la variable independiente y 3 la variable de control.

⁴ Kim J.-O., y Kohout F.: “Multiple Regression Analysis: Subprogram Regression”, en Nie, N., Hull, C. H., Jenkins, J., et. al.: *Statistical Package for the Social Sciences*; McGraw-Hill, Nueva York, 1975, 2ª ed..

⁵ La inclusión de todas las variables mudas resultantes de categorías nominales hace que las ecuaciones normales no puedan ser resueltas, ya que la inclusión de las últimas categorías está completamente determinada por los valores de las primeras categorías ya incluidas en la ecuación.

⁶ Gansner, D., Seegrist, D., Walton G.: “Technique for Defining Subareas for Regional Analysis”, en *Growth and Change*; octubre, 1971.

⁷ Para más detalles, puede verse el capítulo correspondiente a escalas.

⁸ Stouffer, S., et al.: *Measurement and Prediction. Studies in Social Psychology in World War II*, vol. IV; Princeton University.

ÍNDICE

PRÓLOGO. Jorge Padua

I. La organización de un *survey*. Jorge Padua e Ingvar Ahman

Paso I: Orientación en el campo de la investigación y formulación de un sistema de hipótesis

Documentación descriptiva

Estudio de la situación

Documentación explicativa

Resumen

Paso II: La construcción, evaluación y manejo del instrumento para la recolección de datos (cuestionario) y muestreo

Observaciones, entrevistas y cuestionarios

Cuestionario inicial

Muestra para el *pretest*

Entrevistadores para el *pretest*

Organización y evaluación del *pretest*

Diseño de la muestra

Entrenamiento de los entrevistadores

Paso III: La recolección de datos. El trabajo en el campo

Paso IV: El procesamiento de datos

Codificación

Perforación de las tarjetas

Paso V: Análisis

Paso VI: Presentación

Comentarios finales

II. El proceso de investigación. Jorge Padua

Distintos tipos de investigaciones

Teorización en la investigación

Proposiciones e hipótesis

Conceptos, indicadores, índices

Índices

La utilización de puntajes en un indicador, para otro indicador en el que no se tiene respuesta

Apéndice

III. Muestreo. Jorge Padua

Distintos tipos de muestras

A) Muestras probabilísticas

1) El muestreo simple al azar

¿Cómo se decide el tamaño de la muestra?

Selección de una muestra aleatoria en poblaciones de tamaño finito

Ventajas y desventajas del muestreo aleatorio simple

Ejemplo de una muestra aleatoria simple

2) Muestreo sistemático

Ventajas y desventajas

3) Muestreo estratificado

Ejemplo 1: Muestra estratificada proporcional

Ejemplo 2: Muestra estratificada no proporcional

Ventajas y desventajas

4) Muestreo por conglomerados

Ventajas y desventajas

En síntesis

B) Muestras no probabilísticas

1) Las muestras casuales

2) Las muestras intencionales

3) Las muestras por cuotas

C) Muestras para probar hipótesis sustantivas

IV. El cuestionario. Jorge Padua e Ingvar Ahman

Motivación en el entrevistado y el cuestionario como una unidad

El orden de las preguntas
El tamaño del cuestionario
Espacio para las preguntas en el cuestionario
Un ejemplo de la forma de diseñar un cuestionario

Espacio para el empadronamiento general en el cuestionario

Las últimas páginas del cuestionario

Diferentes tipos de preguntas

La pregunta cerrada
La pregunta cerrada “simple”
La pregunta “cerrada” con múltiples respuestas
Técnicas especiales: hojas sueltas
La pregunta abierta
La utilización de preguntas cerradas y abiertas, en conjunto
Ejemplo de un “código múltiple”
Ejemplo de preguntas encadenadas
Preguntas particulares
La pregunta de ingresos
Preguntas sobre ocupación
Preguntas sobre conocimientos

“Lista de control” para cuestionarios y entrevistas

Una “lista de control” para la construcción del cuestionario
Una “lista de control” para el cuestionario final como un todo y para las diversas baterías de preguntas
“Lista de control” de asuntos de entrevistas relacionadas con el cuestionario

Un ejemplo
Cuestionario para jefes de familia

¡Leer esto antes de empezar!

Cuestionario para la familia del alumno

V. La codificación. Héctor Apezechea

La codificación
Distintas clases de preguntas
El código

Nomenclatura

Confección del código

¿Está usted de acuerdo con la actual política del gobierno?

Ejemplos de códigos

Procedimiento de codificación

Codificación que realiza el encuestador

El procedimiento de codificación en la oficina

Instrucciones para la codificación de preguntas abiertas

Procedimientos administrativos. Algunos ejemplos

Procedimiento de revisión y control

Tiempo necesario para la codificación. Errores residuales

Tiempo necesario de la codificación

Errores residuales

Apéndice

- 1) La utilización de símbolos numéricos
- 2) Codificación y cuadros bidimensionales
- 3) Codificación y escalas
- 4) Codificación y tabulación manual
- 5) Codificación y cuestionario

VI. Escalas para la medición de actitudes. Jorge Padua e Ingvar Ahman

Escalas para la medición de actitudes

Actitudes e intereses: algo acerca de su dirección e intensidad

La escala Lickert

- A) La construcción de una escala Lickert
- 1) La construcción de los ítems
- 2) La administración de los ítems a una muestra de jueces

Ejemplo de una versión preliminar

- 3) Asignación de puntajes a los ítems
- 4) Asignación de puntajes totales
- 5) Análisis de los ítems
- 6) La versión final de la escala
- B) Comentarios finales

La escala Thurstone

- A) La construcción de una escala Thurstone

- 1) La construcción de los ítems (versión de los jueces)
- 2) La administración a los jueces

Ejemplo: escala actitud tipo Thurstone

Instrucciones

[Para el ejemplo presentamos solamente los primeros 29 ítems]

- 3) Asignación de valores de escala
- 4) La selección de los ítems
- B) La versión final de la escala
- 5) La adjudicación de puntajes a los sujetos
- C) Comentarios
 - 1) Ítems o series de ítems acumulativos o diferenciados
 - 2) La variación en el instrumento, los sujetos o ambos
 - 3) El problema con las dimensiones
- D) Ventajas y desventajas de la escala Thurstone

La escala Guttman

- A) La construcción de un escalograma Guttman
- 4) El análisis de los ítems
- B) La determinación de los errores
- C) Ejemplos

Referencias

- 1) Alcance de la distribución marginal
- 2) Pauta de errores
- 3) Número de ítems en la escala
- 4) Número de categorías de respuestas
- C) La técnica de la escala discriminatoria de Edwards y Kilpatrick
- D) La técnica H para mejorar la escalabilidad de la escala Guttman
- E) La versión final de la escala
- F) Ventajas y desventajas de la escala Guttman
- G) Comentarios finales
- ¿Variación en el instrumento, en los sujetos, o en ambos?

Método de comparación por pares

- A) Ejemplo de construcción de una escala basada en el método de comparación por pares
 - 1) Versión de los jueces. *Test* de Tipo B
 - 2) Versión de la escala para sujetos. Versión final. *Test* de Tipo A

B) Ejemplo de una escala de comparaciones por pares en un cuestionario (entrevista)

C) Ventajas y desventajas

El diferencial semántico. La escala de Osgood

A) Antecedentes teóricos

B) Dimensiones en el espacio semántico

Dimensiones

C) Construcción

D) Selección de escalas

Análisis de los datos

E) El *test* de Osgood y la medición de actitudes

Instrucciones

El mercado bursátil

Escala de distancia social de Bogardus

A) Procedimientos básicos para la construcción de una escala de distancia social

B) A continuación presentamos los ítems que ilustran esta escala de distancia social

Escala de distancia racial

Instrucciones

C) Flexibilidad de la técnica

Instrucciones

D) Confiabilidad

E) Validez

F) Limitaciones y aplicaciones

Bibliografía recomendada para escalas de medición de actitudes

VII. El trabajo de campo. Ingvar Ahman

1) Descripción del estudio

2) Descripción de la muestra

3) Cómo ponerse en contacto con los respondentes

La presentación

Confidencialidad

Informes a la central de entrevistadores
Instrucciones para la investigación
Verificación en el cuestionario
La inscripción de los *R*

VIII. Análisis de datos: el concepto de propiedad-espacio y la utilización de razones, tasas, proporciones y porcentajes. Carlos Borsotti

A) El concepto de propiedad-espacio

La propiedad-espacio en una variable
La propiedad-espacio en un sistema de dos variables
La propiedad-espacio en un sistema de tres variables
La substrucción

B) Razones, tasas, proporciones y porcentajes

Razones
Tasas
Proporciones
Porcentajes
Medición del cambio porcentual
Bibliografía

IX. Análisis de datos: paquete estadístico para las ciencias sociales (SPSS): oferta y condiciones para su utilización e interpretación de resultados. Jorge Padua

Niveles de medición

- A) Nivel nominal
- B) Nivel ordinal
- C) Nivel intervalar
- D) Nivel por cociente o racional

Programa estadístico del SPSS
Estadística descriptiva

- A) Medidas de tendencia central
Empleo de media, mediana y modo
- B) Medidas de variabilidad o de dispersión

Confiabilidad de los estadísticos
Confiabilidad de diferencia entre estadísticos (subprograma *breakdown*)

- A) Error estándar de la diferencia de medias (Subprograma T-Test)

Tablas de contingencia y medidas de asociación (subprograma *crosstabs*)

Tabulaciones cruzadas

Ji cuadrado (χ^2)

Supuestos y requisitos generales

Coeficiente ϕ (ϕ)

Coeficiente V. de Cramer

Coeficiente de contingencia (C)

El coeficiente Q de Yule

Coeficiente lambda (λ)

Coeficiente τ_b de Goodman y Kruskal

Coeficiente de incertidumbre

Coeficiente tau b

Coeficiente tau c

Coeficiente gamma (γ)

Coeficiente D de Sommer

Coeficiente eta (η)

Correlación biserial (r_b)

Correlación punto-biserial (r_{pb})

Coeficiente de correlación Spearman (ρ)

Coeficiente de correlación tau de Kendall (τ)

Coeficiente de correlación producto-momento de Pearson (r)

Diagrama de dispersión (Scattergram)

Correlación parcial

Análisis de regresiones múltiples (subprograma *regression*)

Ejemplos

Casos especiales en el *path analysis*

Regresiones con variables mudas (*dummy variables*)

Análisis de la varianza unidireccional con variables mudas

Análisis de la varianza multidireccional con variables mudas

Análisis de la varianza y de la covarianza (subprogramas *anova* y *oneway*)

Análisis de la varianza simple

Análisis de la varianza n-dimensional

Análisis factorial (subprograma *factor*)

- A) Preparación
- B) Factorización
- Métodos de factorización en el SPSS
 - a) Factorización principal sin interacción
 - b) Factorización principal con interacción
 - c) Factorización canónica de Rao
 - d) Alfa-factorización
 - e) Imagen-factorización
- C) Rotación
- D) Interpretación

Producto del programa factor

Soluciones terminales para factores rotados ortogonalmente

Soluciones terminales para factores rotados oblicuamente

Análisis discriminante (*subprogram discriminant*)

Algunos ejemplos de análisis discriminante

- a) Ejemplo en dos grupos
- b) Ejemplo con varios grupos

Análisis de escalograma Guttman (*subprograma Guttman scale*)

- a) Alcance de la distribución marginal
- b) Pauta de errores
- c) Número de ítems en la escala
- d) Número de categorías de respuestas

Bibliografía

X. La presentación del informe de investigación. Ingvar Ahman

La lista de control (*check list*)

ibliografía general

En las ciencias sociales es parte esencial de la investigación el acopio de datos mediante el contacto directo o indirecto con informantes. Las técnicas al respecto dependen del problema y del género de interrogantes que la indagación sugiera, del planteamiento teórico general de ésta y de la etapa en que se halle la teoría sustantiva.

El presente manual describe los aspectos operativos de la investigación, entre ellos la estructuración de los cuestionarios para la recolección de datos; además, aborda el proceso de organización de la investigación, el ejercicio de variables y los procedimientos de muestreo y codificación, así como la construcción de escalas, el análisis de la información y las recomendaciones para la presentación del informe de la indagación.

Preparados por cuatro especialistas, todos los capítulos han sido escritos para que puedan ser leídos de forma independiente, pero están ordenados de manera tal que el volumen, en su conjunto, sigue la lógica del proceso de investigación. Así, la obra llena un vacío que desafortunadamente existe en la bibliografía especializada, y resulta de gran valor tanto para quienes estudian o imparten algún curso introductorio de técnica de la investigación como para los investigadores que frecuentemente necesitan consultar los aspectos técnicos de su labor.